



اوپو سسيي نايكولو مي مارا
UNIVERSITI
TEKNOLOGI
MARA



JCRINN

JOURNAL OF COMPUTING RESEARCH AND INNOVATION

VOL. 10 NO. 2 (2025)



[HTTPS://JCRINN.COM/](https://jcrinn.com/)

Journal of Computing Research and Innovation.

Copyright © 2025 UiTM Press.

eISSN: 2600-8793

This work is licensed under the Creative Commons Attribution 4.0 License (CC-BY 4.0)

You are free to:

- Share: copy and redistribute the material in any medium or format.
- Adapt: remix, transform, and build upon the material for any purpose, even commercially.
- Under the following terms:

Attribution: You must give appropriate credit, provide a link to the license, and indicate if changes were made. You may do so in any reasonable manner, but not in any way that suggests the licensor endorses you or your use.

No additional restrictions: You may not apply legal terms or technological measures that legally restrict others from doing anything the license permits.

This work is provided “as is,” and the author/ publisher disclaims any and all warranties, whether expressed or implied, including but not limited to the implied warranties of merchantability and fitness for a particular purpose and non-infringement. The author/ publisher shall not be liable for any damages whatsoever, including without limitation any special, indirect, or consequential damages, arising out of or in connection with the use or inability to use this work.

Journal of Computing Research and Innovation Volume 10, Number 2 (2025).

DOI: 10.24191/jcrinn.v10i2

Associate Editor in Chief : Siti Zulaiha Ahmad (sitizulaiha@uitm.edu.my)

Managing Editor: Mohammad Hafiz Ismail (mohammadhafiz@uitm.edu.my)

Bahagian PJI
Universiti Teknologi MARA,
02600 Arau, Perlis,
MALAYSIA

jcrinn@uitm.edu.my

<https://jcrinn.com>

JOURNAL OF COMPUTING RESEARCH AND INNOVATION

Editorial Board

EDITOR-IN-CHIEF

Shukor Sanim Mohd Fauzi

Universiti Teknologi MARA, Perlis Branch, MALAYSIA

Email: jcrinn@uitm.edu.my

ASSOCIATE EDITOR-IN-CHIEF

Siti Zulaiha Ahmad

Universiti Teknologi MARA, Perlis Branch, MALAYSIA

Email: jcrinn@uitm.edu.my

MANAGING EDITOR

Mohammad Hafiz bin Ismail

Universiti Teknologi MARA, Perlis Branch, MALAYSIA

COPY EDITOR

Nurtihah Mohamed Noor

Universiti Teknologi MARA, Perlis Branch, MALAYSIA

JOURNAL OF COMPUTING RESEARCH AND INNOVATION

Table of Contents

Meteorological Factors Affecting on PM2.5 Concentrations in Bandaraya Melaka

*Nur Alya Adriana Abdullah Sani, Nurkhairany
Amyra Mokhtar*

1-15

COVID-19 Spread in Malaysia Using SIR Model

Nur Nadiah Az-Zahraa, Nurul Najihah Mohamad

16-25

Youth Insights on Factors Influencing Housing Purchase and Rental Decisions

*Mohamad Haizam Mohamed Saraf, Muhammad
Eidlan Hakimi Radzi, Syahmimi Ayuni Ramli,
Mohammad Fitry Md Wadzir*

26-34

Evaluating Students' Bread Preferences in UiTM Perlis: An Application of the Fuzzy Analytical Hierarchy Process

*Khairu Azlan Abd Aziz, Wan Suhana Wan Daud,
Mohd Fazril Izhar Mohd Idris, Nur Aqilah Kamal
Azman, Rizauddin Saian*

35-53

Factory Simulation Construction Method and Implementation of Intelligent Manufacturing

*Bin Zheng, Yan Bai, Soo Siang Yang, Ming Keng
Tan, Jing Song Zhang*

54-68

PetalsPalette: Blossoming Knowledge by Bridging Technology and Nature through Augmented Reality Learning Application for Flower Identification

*Azmatun Farwizah Farid, Nurtihah Mohamed Noor,
Mohd Fitri Yusoff*

69-83

Assessing Knowledge Acquisition and Perceptions in Hands-On Video Editing Workshop for Non-Technical Students

*Zubaidah Bohari, Abdul Hadi Abdul Talip, Satria
Arjuna Julaihi, Rumaizah Che Md Nor, Norizuandi
Ibrahim*

84-95

**Analysing the Potential Vulnerabilities in
Online Voting Protocols: Homomorphic
Encryption Approach**

*Charles Yao Azameti, William Asiedu, George
Asante*

..... 96-110

**Determinants of Halal Food Purchasing
Decisions Among Undergraduates at UiTM
Tapah**

*Ilya Zulaikha Zulkifli, Nurul Husna Jamian, Samsiah
Abdul Razak, Ahmad Nur Azam Ahmad Ridzuan*

..... 111-118

**Cyberpark: An-IoT based Automated
Parking Management and Monitoring
System using RFID and ESP32**

*Nur Amalina Muhamad, Norhalida Othman, Nor
Diyana Md Sin, Noor Hasliza Abdul Rahman,
Muhammad Iskandar Mohd Rodzi*

..... 119-129

**A Web-Based System for Managing Student
Attendance and Assessment Submissions: An
SDLC Waterfall Model Approach**

*Siti Balqis Mahlan, Jamal Othman, Maisurah
Shamsuddin, Syarifah Adilah Mohamed Yusoff,
Zalina Othman*

..... 130-143

**Hydrological Risk Assessment of the Johor
River: Flood Frequency Analysis with L-
Moments**

*Nur Hanida Othman, Basri Badyalina, Sheril Haiza
Shah Izan, Ani Shabri, Muhammad Fadhil Marsani,
Fatin Farazh Ya'acob*

..... 144-155

**User Acceptance Testing of the ‘Furwish’
Location-Based Mobile Application: An
Empirical Study on Pet Adoption and
Surrender Services**

*Marnie Umirah Mazlan, Muhammad Nabil Fikri
Jamaluddin, Iman Hazwam Abd Halim, Alif Faisal
Ibrahim, Mohd Faris Mohd Fuzi*

..... 156-169

**Impact of Blackhole and Wormhole Attacks
on DSDV Routing Protocol in VANET:
Behavioural Analysis**

*Ahmad Yusri Dak, Nuramarina Nasruddin, Nur
Khairani Kamaruddin*

..... 170-181

AI-Driven Forgery Detection in Offline Handwriting Signatures: Advances, Challenges, and the Role of Generative Adversarial Networks

Safura Adeela Sukiman, Nor Azura Husin, Hazlina Hamdan, Masrah Azrifah Azmi Murad

182-197

CNN Integrated Mobile Application: Food Image Recognition for Recipe Generation

Muhammad Imran Nor Azlan Shah, Norlina Mohd Sabri, Gloria Jennis Tan, Zhiping Zhang

198-215

Enhancing Pineapple Cultivar Classification: A Framework for Image Quality, Feature Extraction, and Algorithmic Refinement

Mohamad Faizal Ab Jabal, Muhammad Irfan Rozlan, Azrina Suhaimi, Harshida Hasmy

216-229

Thyroid Insight: Navigating Disease Data Through Interactive Visualization with Prediction

Mohd Nizam Osman, Azim Md Nasib, Khairul Anwar Sedek, Nor Arzami Othman, Mushahadah Maghribi

230-241

Heart Failure Detection Using Scaled Conjugate Gradient Method and Naïve Bayes

Norpah Mahat, Norazwana Saidin @ Zubir, Jasmani Bidin, Mohamad Najib Mohamad Fadzil, Sharifah Fhahriyah Syed Abas, Siti Sarah Raseli

242-254

Enhancing Photovoltaic Panel Inspection using RGB Image-Based Detection and Image Processing Techniques

Suhaili Beeran Kutty, Mohamad Ad-Fadhil Musa, Murizah Kassim, Puteri Nor Ashikin Megat Yunus

255-265

Bibliometric Analysis of Research on Firth Penalized Logistic Regression in Addressing Complete Separation

Nurul Husna Jamian, Ahmad Zia Ul-Saufie, Mohammad Nasir Abdullah

266-280

**Navigating AI in Higher Education:
Examining Over-Reliance and Plagiarism
among UiTM Tapah Students**

*Nor Aslily Sarkam, Nor Hazlina Mohammad,
Nurizatie Farhanim Saiful Lizam, Nurul Ain Azmi,
Nadhira Yasmin Mazli, Mizan Qamalia Mohd
Dzahir, Ros Amira Arisha Md Isa*

..... 281-294

**Exploring the Effect of Shape Parameters on
Font Design Using Rational Quadratic
Trigonometric Curves**

*Noor Khairiah Razali, Nur Ain Fitrah Azhari, Siti
Musliha Nor-Al-Din, Nursyazni Mohamad Sukri*

..... 295-305

**Key Determinants of Athletic Success: A
Fuzzy AHP Approach**

*Khairu Azlan Abd Aziz, Wan Suhana Wan Daud,
Mohd Fazril Izhar Mohd Idris, Muhammad Amsyar
Izmi, Rizauddin Saian*

..... 306-318

**A Fuzzy Analytic Hierarchy Process (FAHP)
Approaches for Investment Decision-Making
in Malaysia**

*Mohd Fazril Izhar Mohd Idris, Khairu Azlan Abd
Aziz, Ardini Athirah Mhd Munawar*

..... 319-334

**Evaluating User Acceptance of the UnityHub
Web-based System in Higher Education: A
Descriptive Analysis**

*Norazah Umar, Nurhafizah Ahmad, Jamal Othman,
Rozita Kadar*

..... 335-348

Meteorological Factors Affecting on PM_{2.5} Concentrations in Bandaraya Melaka

Nur Alya Adriana Abdullah Sani¹, Nurkhairany Amyra Mokhtar^{2*}

¹Statistics Studies, Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Negeri Sembilan Branch, Seremban Campus, 70300 Seremban, Negeri Sembilan, Malaysia.

²Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Johor Branch, Segamat Campus, 85000 Segamat, Johor, Malaysia.

ARTICLE INFO

Article history:

Received 12 March 2025

Revised 26 March 2025

Accepted 17 April 2025

Published 1 September 2024

Keywords:

Fine Particulate Matter
Meteorological Factors
Wind Speed
Ambient Temperature
Relative Humidity
Bandaraya Melaka

DOI:

10.24191/jcrinn.v10i2.514

ABSTRACT

The most hazardous pollutant, PM_{2.5} has caused serious environmental and public health issues. Living in an area with PM_{2.5} pollution can lead to respiratory problems as it can reach the human bloodstream through inhalation. Bandaraya Melaka was chosen to be the study area as it experienced urban tropical environments, where meteorological parameters, including ambient temperature, wind speed, and relative humidity, substantially impact the concentration of PM_{2.5}. Hence, this study investigates the relationship between the concentration of PM_{2.5} and these meteorological factors using daily data collected from January to early July 2019. Statistical analyses were conducted after model adequacy checking on the linearity, normality, homoscedasticity, independence of the error term, no multicollinearity, and the absence of an outlier in multiple linear regression fulfilling through five iterations of outlier removal. The F-test revealed a relationship exists between meteorological factors and the concentration of PM_{2.5}. The results indicate a moderate positive relationship between meteorological factors and the concentration of PM_{2.5}, with only 38% of the total variation in the concentration of PM_{2.5} explained by these factors. In the t-test, all meteorological variables were found to significantly influence the concentration of PM_{2.5}, and the final model, containing all factors, was identified as the best-fitting model supported with the lowest AIC and BIC values. This study contributed to better insight into forecasting the quality of the air in tropical urban environments, addressing a research gap in Bandaraya Melaka. These findings are essential for designing effective air quality management strategies, protecting public health, and supporting urban development.

1. INTRODUCTION

Fine particulate matter (PM_{2.5}) was among the most hazardous air contaminants. PM_{2.5} is an often-examined particulate matter measure, due to its extensive array of acknowledged detrimental health impacts,

^{2*} Corresponding author. E-mail address: nurkhairany@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.514>

especially with cardiovascular health concerns (Krittanawong et al., 2023). It comprised minuscule solid particles and liquid droplets with a diameter of $2.5 \mu\text{g}/\text{m}^3$ (Wang & Ogawa, 2015). The size made it significantly smaller than numerous typical particles, including human hair. $\text{PM}_{2.5}$ was approximately 20 times smaller than the diameter of a single hair strand (U.S. Environmental Protection Agency, 2023). It was frequently referred to as dust, soot, dirt, smoke, or aerosols because its tiny size made it visible only under an electron microscope (IQAir, 2022).

Many sources can influence the concentration of $\text{PM}_{2.5}$ in the atmosphere. One of the sources was meteorological factors, which were among the significant factors. Kayes et al. (2019) mentioned that meteorological factors have a crucial role in determining the air quality in the atmosphere. Kayes et al. (2019) also mentioned in the study focusing on urban regions that meteorological parameters, specifically ambient temperature, relative humidity, and wind speed, were important contributors to the fluctuation of $\text{PM}_{2.5}$ which was the reason this study chose the parameters.

In addition, meteorological factors were chosen in addition to the topography and climate of Bandaraya Melaka, which demonstrated that these meteorological parameters and Bandaraya Melaka were correlated. Unlike other areas, Bandaraya Melaka is located on Peninsular Malaysia's west coast, where the wind from the coast is strong enough to cause an effect on the level of $\text{PM}_{2.5}$ (Lim, 2015). In addition, Bandaraya Melaka was selected as it is in a tropical region with high temperature, humidity, and seasonal winds, which could be influential in $\text{PM}_{2.5}$ concentration (Latif et al., 2018).

Despite assumptions that urban areas contribute less to fine particulate pollution compared to industrial zones, research by Gao et al. (2018) and Y. Wang et al. (2016) challenges this notion, identifying urban centers as significant contributors to $\text{PM}_{2.5}$ pollution. The lack of focused research on meteorological factors in urban tropical environments creates a knowledge gap that this study aims to address. By analysing daily data from January to early July 2019 using multiple linear regression, this research seeks to examine the relationship between ambient temperature, relative humidity, wind speed, and $\text{PM}_{2.5}$ concentrations in Bandaraya Melaka.

Wind speed affects $\text{PM}_{2.5}$ concentrations as it impacts the dispersion of airborne pollutants. Chen et al. (2020), in China, when studying the relationship between meteorological parameters and the concentration of fine particulate matter, detected a negative association between the concentration of $\text{PM}_{2.5}$ and wind speed. They explained that stronger winds have a direct blowing-off force that further promotes the dispersion of particles.

Ambient temperature could influence the phenomenon of temperature inversions in China that block the pollutants from going upward, resulting in excessive concentrations of fine particulate matter (Li et al., 2017). Ma et al. (2021) studies on the $\text{PM}_{2.5}$ and O_3 behaviour to meteorological factors in China found that higher temperatures facilitate chemical reactions between vehicle emissions and industrial activities, forming fine particulate matter resulting in high $\text{PM}_{2.5}$ levels (Majid et al., 2020). A study by Amil et al. (2016) on the variability of $\text{PM}_{2.5}$ in Klang Valley from August 2011 to July 2012 suggests that higher temperatures might form dry conditions that suppress air circulation, trapping pollutants near the surface. The reduced vertical mixing leads to the buildup of $\text{PM}_{2.5}$ concentrations at the surface.

How and Ling (2016) explore at Universiti Teknologi Malaysia (UTM) at Skudai, Johor, a region with few rainfalls a year due to high humidity for the rainy season period. Through correlation evaluation, their results showed that temperature and relative humidity aid the aggregation and subsequent removal of particulate matter via precipitation, helping to lower airborne $\text{PM}_{2.5}$ concentrations. Studies looking at data from China from 2006 to 2014 found that high relative humidity aided chemical reactions. This process converts sulfur dioxide to sulfate aerosols and consequently enhances the concentration of aerosolised $\text{PM}_{2.5}$ in the air (Fang et al., 2017). Ramli et al. (2024) and Dahari et al. (2020) used Pearson Correlation Analysis to examine meteorological influences on air pollutants in Malaysia across seasons. Ramli et al. (2024) found that NO_2 , CO, and O_3 correlated with climatic factors, with stronger associations during the

Southwest monsoon. Dahari et al. (2020) reported a positive correlation between temperature and PM_{2.5} in Skudai, Johor, and a negative correlation with relative humidity and wind speed, with the strongest effects during the Southwest monsoon. Both studies highlight the role of meteorological factors in air pollution variability.

2. METHODOLOGY

The data was collected from 1st January 2019 to 12th July 2019, a total of 193 observations from the Department of Environment (DOE) and the Malaysia Meteorological Department (MET Malaysia). The study was based on PM_{2.5} concentration data from the Department of Environment. Data from an air monitoring station in Bandaraya Melaka in 2019. Met Malaysia provided meteorological data during the same period, such as ambient temperature, relative humidity, and wind speed. This study analysed the dependent variable of PM_{2.5} concentration with independent factors of wind speed, relative humidity, and ambient temperature,

Multiple linear regression was employed in this study because there was more than one predictor variable (Montgomery et al., 2021). This method was used to analyse the relationship between the concentration of PM_{2.5} and meteorological factors, including wind speed, ambient temperature, and relative humidity. Before undergoing multiple linear regression, model adequacy checking should be done. The normality assumption is checked by using the Kolmogorov-Smirnov test, the linearity is assessed through the scatter plots, the homoscedasticity assumption is checked through the scatter plot of residuals against the predicted values, the independence of error term assumption is evaluated through the scatter plot of residuals against date, the multicollinearity assumption is checked through the values of Tolerance (TOL) and Variance Inflation Factor (VIF), and the outliers are detected through the scatter plot of studentized residuals against predicted values. Then, the validity of the test, coefficient of determination, t-test and AIC-BIC were checked to ensure that the model that has a significant relationship is the best-fit model. Data was analysed using SPSS and R-studio.

3. FINDINGS AND DISCUSSION

3.1 Set of variables

This study examined the dependent variable of PM_{2.5} concentration with independent factors of wind speed, relative humidity, and ambient temperature to analyse the relationship between PM_{2.5} concentration with three meteorological factors, which are wind speed, relative humidity, and ambient temperature as shown in Table 1.

Table 1. Set of variables

Variable Name	Variable Type	Description
PM2.5 Concentration	Numerical	The reading on the concentration of PM2.5 (µg/m ³)
Relative Humidity	Numerical	Percentage of Relative Humidity (%)
Wind Speed	Numerical	Speed of the Wind (ms ⁻¹)
Ambient Temperature	Numerical	Atmospheric Temperature (°C)

3.2 Model equation

The model equation can be expressed as follows:

$$PM_{2.5} = \widehat{\beta}_0 + \widehat{\beta}_1 x_1 + \widehat{\beta}_2 x_2 + \widehat{\beta}_3 x_3 + \epsilon$$

where,

y_i : $PM_{2.5}$ concentration

x_1 : Relative Humidity

x_2 : Wind Speed

x_3 : Ambient Temperature

$\widehat{\beta}_0$: the intercept

$\widehat{\beta}_1, \widehat{\beta}_2, \widehat{\beta}_3$: The regression coefficients of the independent variables

ϵ : the error term.

3.3 Model adequacy checking

As per Montgomery et al. (2021), it was essential to go through model adequacy checking before proceeding with the regression analysis since it may help throw light on an error that could cause the analysis not to match the actual scenario and become unreliable.

3.3.1 Before removing outliers

Normality Assumption:

Table 2 shows the result of the Kolmogorov-Smirnov test. The p-value obtained is 0.001, smaller than 0.05 (α). This indicates a significant deviation from the normal distribution, which can be concluded that the normality assumption is not satisfied.

Table 2. Kolmogorov-Smirnov test before the outliers were removed

Kolmogorov - Smirnov	
Statistic	Significance
0.284	0.001

Linearity Assumption:

Fig. 1 shows the scatterplot matrix of $PM_{2.5}$ concentration against ambient temperature, wind speed, and relative humidity. The pattern of the scatterplot appeared to have a randomly scattered point with an increasing pattern between ambient temperature and $PM_{2.5}$ concentration. Other than that, the relative humidity and $PM_{2.5}$ concentration had shown a decreasing pattern with a randomly scattered point. Next, wind speed has a negative relationship with $PM_{2.5}$ concentration as it appears to have a randomly scattered

plot with a decreasing pattern. Since, the scatterplot between $PM_{2.5}$ concentration and meteorological parameters shows a randomly scattered point, it can be concluded that the linearity assumption is satisfied.

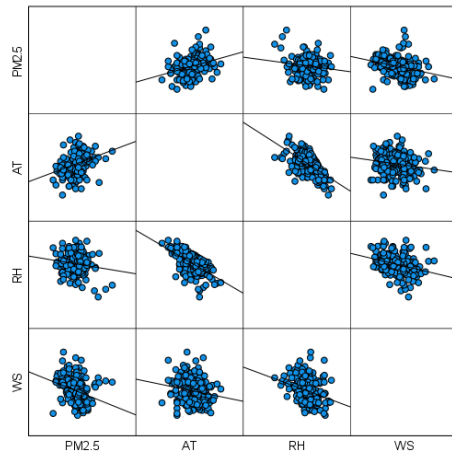


Fig. 1. Scatterplot matrix before the outlier were removed

Homoscedasticity Assumption:

Fig. 2 shows the scatter plot of residuals against the predicted values. The distribution of residuals appears to be randomly distributed, suggesting that the assumption of constant variance is fulfilled.

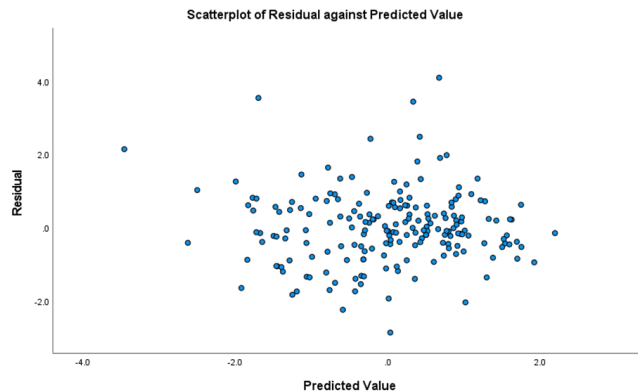


Fig. 2. Scatter plot of residual against predicted value before outliers were removed

Independence of Error Term Assumption:

Fig. 3 shows the scatter plot of residuals against date, illustrating a cyclical pattern. This indicates the presence of autocorrelation, which may violate the assumption of the independence of the error term.

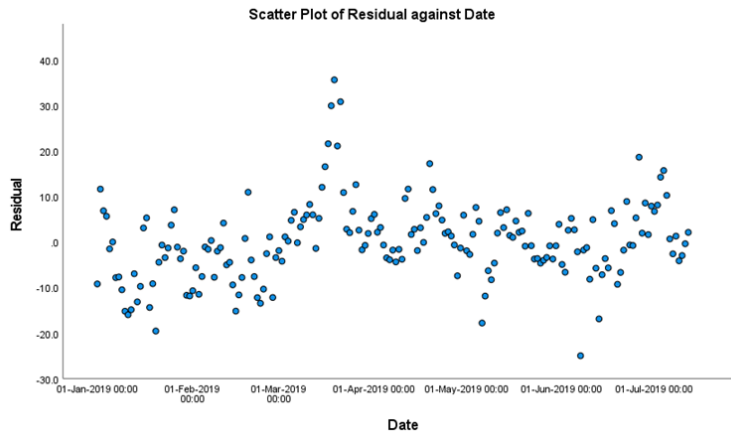


Fig. 3. Scatter plot of residual against date before outliers were removed

Multicollinearity Assumption:

Table 3 shows the values of Tolerance (TOL) and Variance Inflation Factor (VIF). The tolerance values obtained are 0.485 for ambient temperature, 0.456 for relative humidity, and 0.709 for wind speed. All the meteorological factors have TOL values greater than 0.2, suggesting that multicollinearity does not exist. Meanwhile, the VIF values are 2.060 for ambient temperature, 2.193 for relative humidity, and 1.410 for wind speed, which were below 10, further confirming that there is no multicollinearity exists between them.

Table 3. Multicollinearity test

Variables	Tolerance	VIF	Presence of Multicollinearity
Ambient Temperature (AT)	0.485	2.060	No
Relative Humidity (RH)	0.456	2.193	No
Wind Speed (WS)	0.709	1.410	No

Presence of Outlier Assumption:

Fig. 4 shows the scatter plot of studentized residuals against predicted values. Three points are beyond the cutoff point (± 3), which indicates the presence of outliers in the concentration of $PM_{2.5}$ (Montgomery et al., 2021). These outliers may excessively influence the regression model, which can reduce its predictive accuracy.

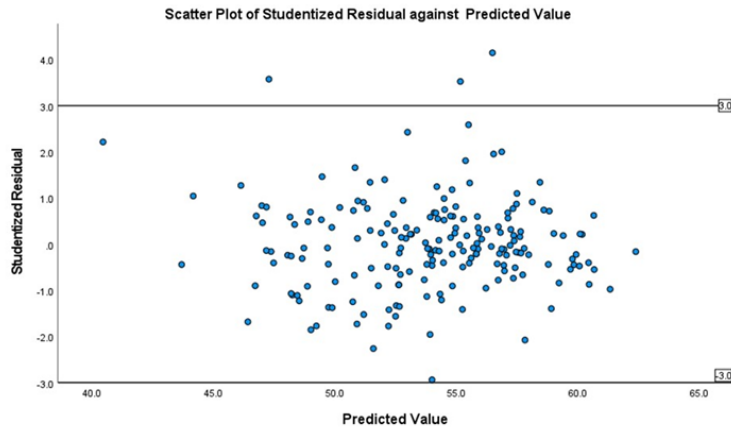


Fig. 4. Scatter plot of studentized residual by predicted value before the outliers were removed

Fig. 5 shows the scatter plot of studentized residuals against leverage values. This plot indicates the presence of outliers, as some points exceed the cutoff of $\frac{2(4)}{193} = 0.04145$. This suggests that the assumption of the regression model is violated due to the presence of outliers in the meteorological factors (x).

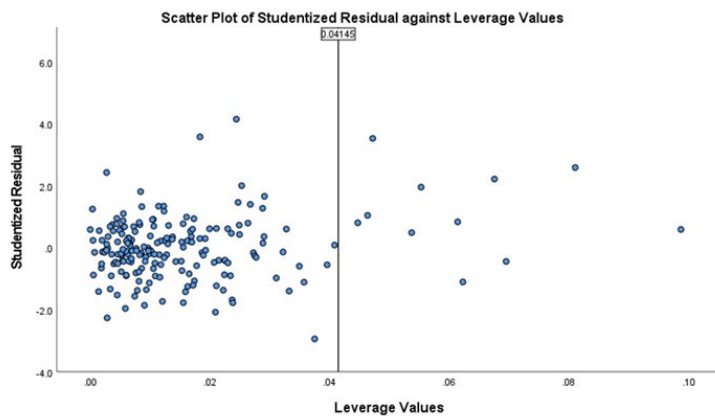


Fig. 5. Scatter plot of studentized residual by leverage before the outliers were removed

3.3.2 Outlier detection

Based on Table 4, after analysing the data, 24 observations are identified as influential points, which require approximately five iterations to be excluded from the analysis.

Table 4. Removed observation detected as outliers

Days	Cook's Distance	DFFITS	DFB0	DFB1	DFB2	DFB3
1	0.0220	-	-	0.2461	-	-
26	-	0.4993	7.7319	-	-	-
27	-	-	3.2687	-	-	-
42	-	-	3.6710	-	-	-
43	-	-	1.2218	-	-	-
45	-	-	3.1235	-	-	-
50	0.0354	0.5785	-	-	-	-
60	0.0363	0.5804	-	-	-	-
73	0.0377	0.6439	7.3495	-	-	-
74	0.1313	1.1447	22.2615	-	-	-
75	0.0614	1.0591	14.6419	-	-	-
76	0.1578	2.0278	22.9148	-	-	-
77	0.1717	1.6490	23.2676	-	-	-
78	0.1311	1.0834	7.4278	-	-	-
79	0.0247	-	-	-	-	-
80	0.0770	0.7340	8.4760	-	-	-
130	-	-	1.9026	-	-	-
158	0.1795	-	-	0.4897	-	0.2057
162	-	0.5606	-	0.1579	-	-
163	0.1288	1.1251	22.1357	-	-	-
166	-	-	2.4425	-	-	-
168	-	0.3574	-	-	-	-
173	-	0.4793	-	-	-	-
177	0.0961	1.4538	-	-	-	0.1664

3.3.3 After removing outlier

Normality Assumption

Table 5 shows the results of the Kolmogorov-Smirnov test after the outliers were removed. The p-value is 0.2, which is greater than 0.05 (α), indicating that the normality assumption is satisfied.

Table 5. Kolmogorov - Smirnov test after outliers were removed

Kolmogorov - Smirnov	
Statistic	Significance
0.045	0.2

Linearity Assumption:

Based on Fig. 6, the scatter plot matrix of PM_{2.5} against ambient temperature, wind speed, and relative humidity. The scatter plot between the concentration of PM_{2.5} and the ambient temperature appears to have a randomly scattered point with an increasing pattern. Next, the scatterplot of PM_{2.5} against relative humidity shows an increasing pattern with a randomly scattered point. Then, the scatter plot of PM_{2.5} concentration against the wind speed shows a decreasing pattern. Since, the scatterplot between PM_{2.5} concentration and meteorological parameters shows a randomly scattered point, it can be concluded that the linearity assumption is satisfied.

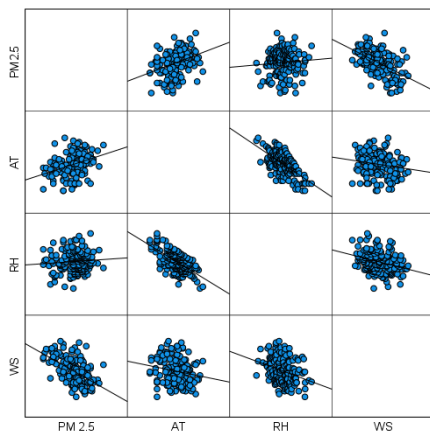


Fig. 6. Scatter plot matrix after the outliers were removed

Homoscedasticity Assumption:

Fig. 7 shows the scatter plot of residuals against predicted values, representing the assumption of homoscedasticity. After removing the outlier, the scatter plot displays a more even distribution of residuals across the predicted values, indicating that the homoscedasticity assumption is satisfied.

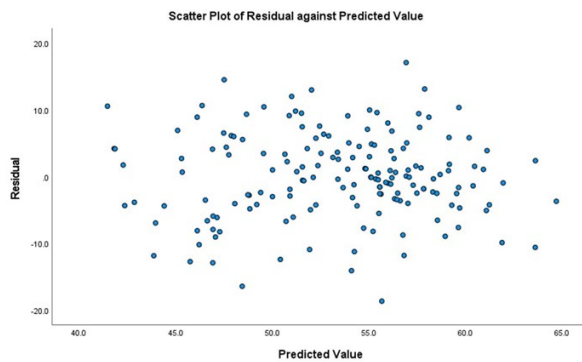


Fig.7. Scatter plot of residual against predicted value after the outliers were removed

Independence of an Error Term Assumption:

Fig. 8 presents the scatter plot of residuals against the dates after removing the outlier. The plot appears to be scattered randomly. This indicates that the independence of the error term is satisfied.

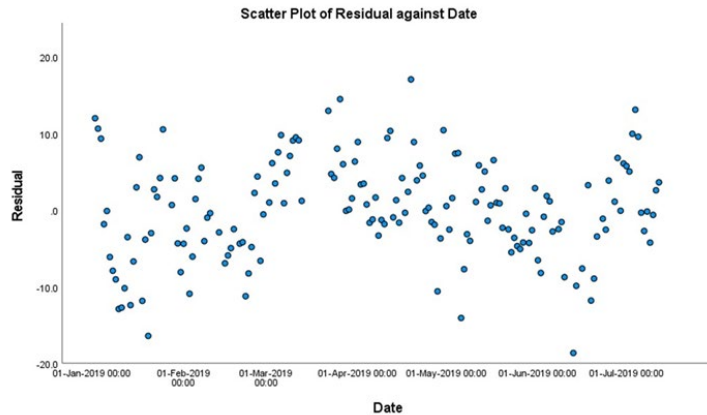


Fig. 8. Scatter plot of residual by date after the outliers were removed

Multicollinearity Assumption:

Table 6 shows the values of Tolerance (TOL) and Variance Inflation Factor (VIF). The tolerance values obtained are 0.429 for ambient temperature, 0.401 for relative humidity, and 0.663 for wind speed. All the meteorological factors have TOL values greater than 0.2, suggesting that multicollinearity does not exist. Meanwhile, the VIF values are 2.333 for ambient temperature, 2.495 for relative humidity, and 1.509 for wind speed, which were below 10, further confirming that no multicollinearity exists between them.

Table 6. Multicollinearity among the variables after the outliers were removed

Variables	Tolerance	VIF	Presence of Multicollinearity
Ambient Temperature	0.429	2.333	No
Relative Humidity	0.401	2.495	No
Wind Speed	0.663	1.509	No

Scatterplot of Studentized Residual against Predicted Value:

Based on Fig. 9 shows the scatter plot of studentized residuals against predicted values. The plot shows that no points exceed the cutoff point for y , indicating that the absence of outliers in y is satisfied.

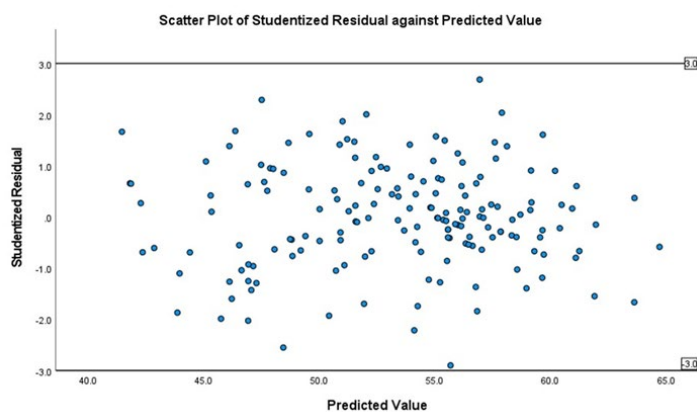


Fig. 9. Scatter plot of studentized residual by predicted value after the outliers were removed

Scatter Plot of Studentized Residual against Leverage Value:

Fig. 10 shows the scatter plot of studentized residuals against leverage values. The plot identifies approximately four points that exceed the cutoff point, which is $\frac{2(4)}{169} = 0.04734$. However, further analysis confirms that these points are not influential and do not significantly affect the regression results. Consequently, no additional observations are removed at this stage, and there are no remaining outliers in this iteration. Hence, the assumption about the absence of outliers is fulfilled.

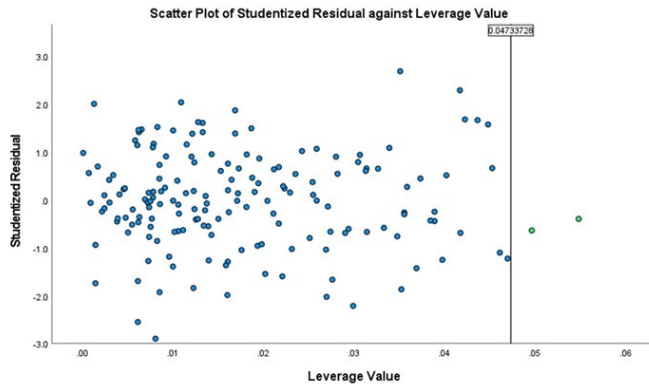


Fig. 10. Scatter plot of studentized residual by leverage value after the outliers were removed

3.4 Modelling data

3.4.1 Testing the validity of the model (F-Test)

Based on the results in Table 7, the value of the F test is 33.749 from the division of MSR and MSE. The value obtained is higher than the critical F value, which is 2.60, and the p-value of 0.001 is lower than the significance level of 0.05. Both the F test and the p-value results indicate that there is a significant relationship between the concentration of $PM_{2.5}$ and meteorological factors, which are wind speed, relative humidity and ambient temperature.

Table 7. F-Test of the model

MSR	MSE	F	$F_{3,165,0.05}$	p-value
1420.294	42.084	33.749	2.60	0.001

3.4.2 Coefficients of determination (R^2)

Table 8 illustrates the association between meteorological factors and the concentration of $PM_{2.5}$. The strength of the relationship is 0.617, which indicates a moderate positive relationship. Then, the coefficient of determination obtained is 0.380. This suggests that approximately 38% of the total variation in the concentration of $PM_{2.5}$ is explained by meteorological factors, while the remaining 62% is attributed to other factors.

Table 8. Coefficient of determination

R	Strength	R-Square
0.617	Moderate Positive	0.380

3.4.3 Testing the significance of the model (T-test)

Based on Table 9, the p-values for ambient temperature, relative humidity, and wind speed are 0.001, 0.011, and 0.001. All meteorological factors have a p-value less than the significance level of 0.05. This indicates that all three variables are significant predictors of $PM_{2.5}$ concentration.

Table 9. T-test of the model

Model	t	Significance
Constant	-2.480	0.014
Ambient Temperature	4.749	0.001
Relative Humidity	2.561	0.011
Wind Speed	-5.051	0.001

3.4.4 The Akaike Information Criterion (AIC) and Bayesian Information Criterion (BIC)

Table 10 presents the AIC and BIC values for each model. As shown, Model 2 has the highest AIC and BIC values, while Model 7 has the lowest. According to Kimura and Waki (2016), the model with the smallest AIC and BIC values is considered the best among the candidates. Therefore, the best model is Model 7, which includes a combination of all predictor variables which is $PM_{2.5} = -104.264 + 4.821AT + 0.470RH - 1.793WS$.

Table 10. AIC & BIC for choosing the best model

Model	Predictors	Model Equation	AIC	BIC
Model 1	AT only	$-104.264 + 4.821AT$	1171.933	1181.323
Model 2	RH only	$-104.264 + 0.470RH$	1193.369	1202.758
Model 3	WS only	$-104.264 - 1.793WS$	1137.171	1146.560
Model 4	AT + WS	$-104.264 + 4.821AT - 1.793WS$	1122.144	1134.663
Model 5	AT + RH	$-104.264 + 4.821AT + 0.470RH$	1139.851	1152.371
Model 6	WS + RH	$-104.264 + 0.470RH - 1.793WS$	1137.210	1149.730
Model 7	AT + WS + RH	$-104.264 + 4.821AT + 0.470RH - 1.793WS$	1117.556	1133.205

3.4.5 Model equation

$$PM_{2.5} = -104.264 + 0.470 * \text{Relative Humidity} - 1.793 * \text{Wind Speed} + 4.821 * \text{Ambient Temperature} \quad (1)$$

Regarding the regression coefficient of constant, it can be said that on average, when the meteorological factors are zero, the reading of $PM_{2.5}$ concentration will be $-104.264\mu g m^{-3}$. While the coefficient of relative humidity represents that with all other variables remaining constant, on average, when relative humidity increases by 1%, the reading of $PM_{2.5}$ concentration increases by $0.470\mu g m^{-3}$. Although the coefficient of wind speed represents that with all other variables held constant, on average, when wind speed increases by $1ms^{-1}$, the reading of $PM_{2.5}$ concentration decreases by $1.793\mu g m^{-3}$. While the coefficient of ambient temperature represents that with all other variables remaining constant, on average, when the ambient temperature increases by $1^{\circ}C$, the reading of $PM_{2.5}$ concentration increases by $4.821\mu g m^{-3}$.

3.5 Discussion

This study used multiple linear regression to examine the impact of meteorological factors on PM_{2.5} concentrations in Bandaraya Melaka. The F-test confirmed the model's statistical significance ($p = 0.001$), indicating that ambient temperature, relative humidity, and wind speed influence PM_{2.5} levels. The coefficient of determination ($R^2 = 0.38$) indicates that 38% of the total variation in PM_{2.5} is explained by meteorological factors. This aligns with Wise and Comrie (2005), who reported low R^2 values (12% to 49%) in similar studies, suggesting that other factors, such as emissions and transboundary pollutants, could be the main contributor to the total variation in PM_{2.5}. Although meteorological factors are not the primary contributor to PM_{2.5} levels, their role in facilitating pollutant transport, transformation, and deposition underscores their relevance in this study.

The t-test showed all parameters were significant: temperature ($p = 0.001$, $r = 0.353$) and humidity ($p = 0.011$, $r = 0.079$) had weak positive relationships with PM_{2.5}, while wind speed ($p = 0.001$, $r = -0.536$) had a moderate negative relationship. These findings are consistent with previous studies attributing PM_{2.5} increases to photochemical reactions and the presence of moisture causing accumulation of the pollutants, while atmospheric mixing processes (both horizontal advection and vertical convection) facilitate the dispersion of pollutants to another area (How & Ling, 2016; Dahari et al., 2020).

AIC and BIC confirmed that the full model, including all meteorological parameters, provided the best fit. The final model suggests that temperature and humidity positively influence PM_{2.5}, while wind speed has a negative effect. These findings highlight the role of meteorological factors in air quality but also emphasize the need to consider non-meteorological contributors. Other factors such as emissions, human activities may be considered in future research.

4. CONCLUSION

This study investigates the relationship between PM_{2.5} concentration and meteorological factors in Bandaraya Melaka using multiple linear regression. The results confirm that ambient temperature, relative humidity, and wind speed significantly impact PM_{2.5} levels, collectively explaining 38% of its variation. While meteorological factors are not the primary contributors, they still play a role in PM_{2.5} fluctuation. The final model suggests that temperature and humidity positively influence PM_{2.5}, whereas wind speed has a negative effect. Though this model may not be ideal for direct prediction, it provides valuable insights into the interaction between meteorology and air pollution. Additionally, this study highlights the importance of meteorological factors in mitigating air pollution, particularly in scenarios such as transboundary haze. Understanding how these parameters influence PM_{2.5} can help authorities implement preventive measures, such as cloud seeding, to reduce pollution levels. Moreover, public awareness of these relationships can aid in environmental planning and air quality management.

To enhance future research, several aspects should be considered. Expanding the study to different regions of Malaysia would provide comparative insights into PM_{2.5} behaviour and help identify location-specific pollution sources. Incorporating additional meteorological variables such as rainfall, atmospheric pressure, and wind direction could improve model accuracy and offer a more comprehensive understanding of PM_{2.5} fluctuations. Furthermore, extending the study duration to 12 months or multiple years would allow for a better understanding of seasonal variations and long-term trends in PM_{2.5} levels. By addressing these aspects, future studies can provide stronger evidence for air pollution management strategies and contribute to more effective environmental policies.

5. ACKNOWLEDGEMENTS/FUNDING

The authors would like to acknowledge the support of Universiti Teknologi Mara (UiTM), Cawangan Negeri Sembilan, for providing the facilities and financial support for this research.

6. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

7. AUTHORS' CONTRIBUTIONS

Nur Alya Adriana Abdullah Sani: Conceptualisation, methodology, formal analysis, investigation, and writing- original draft; **Nurkhairany Amyra Mokhtar:** Conceptualisation, supervision, writing- review and editing, and validation.

8. REFERENCES

- Amil, N., Latif, M. T., Khan, M. F., & Mohamad, M. (2016). Seasonal variability of PM_{2.5} composition and sources in the Klang Valley urban-industrial environment. *Atmospheric Chemistry and Physics*, 16(8), 5357–5381. <https://doi.org/10.5194/acp-16-5357-2016>
- Chen, Z., Chen, D., Zhao, C., Kwan, M.-p., Cai, J., Zhuang, Y., Zhao, B., Wang, X., Chen, B., Yang, J., et al. (2020). Influence of meteorological conditions on PM_{2.5} concentrations across China: A review of methodology and mechanism. *Environment international*, 139, 105558. <https://doi.org/10.1016/j.envint.2020.105558>
- Dahari, N., Latif, M. T., Muda, K., Norelyza, N., et al. (2020). Influence of meteorological variables on suburban atmospheric PM_{2.5} in the southern region of Peninsular Malaysia. *Aerosol and Air Quality Research*, 20(1), 14–25. <https://doi.org/10.4209/aaqr.2019.06.0313>
- Fang, C., Zhang, Z., Jin, M., Zou, P., Wang, J., et al. (2017). Pollution characteristics of PM_{2.5} aerosol during haze periods in Changchun, China. *Aerosol and Air Quality Research*, 17(4), 888–895. <https://doi.org/10.4209/aaqr.2016.09.0407>
- Gao, J., Wang, K., Wang, Y., Liu, S., Zhu, C., Hao, J., Liu, H., Hua, S., & Tian, H. (2018). Temporal-spatial characteristics and source apportionment of PM_{2.5} as well as its associated chemical species in the Beijing-Tianjin-Hebei region of China. *Environmental pollution*, 233, 714–724. <https://doi.org/10.1016/j.envpol.2017.10.123>
- How, C. Y., & Ling, Y. E. (2016). The influence of PM_{2.5} and PM₁₀ on air pollution index (API). *Environmental Engineering, Hydraulics and Hydrology: Proceeding of Civil Engineering, Universiti Teknologi Malaysia, Johor, Malaysia*, 3, 132.
- IQAir. (2022). Iqair — First in air quality. <https://www.iqair.com/newsroom/pm2-5>
- Kayes, I., Shahriar, S. A., Hasan, K., Akhter, M., Kabir, M., & Salam, M. (2019). The relationships between meteorological parameters and air pollutants in an urban environment. *Global Journal of Environmental Science and Management*, 5(3), 265–278.
- Kimura, K., & Waki, H. (2016). Minimization of Akaike's information criterion in linear regression <https://doi.org/10.24191/jcrinn.v10i1>

- analysis via mixed integer nonlinear program. *Optimization Methods and Software*, 33, 633–649. <https://doi.org/10.1080/10556788.2017.1333611>
- Krittanawong, C., Qadeer, Y. K., Hayes, R. B., Wang, Z., Virani, S., Thurston, G. D., & Lavie, C. J. (2023). PM_{2.5} and cardiovascular health risks. *Current problems in cardiology*, 48(6), 101670. <https://doi.org/10.1016/j.cpcardiol.2023.101670>
- Latif, M. T., Othman, M., Idris, N., Juneng, L., Abdullah, A. M., Hamzah, W. P., ... & Jaafar, A. B. (2018). Impact of regional haze towards air quality in Malaysia: A review. *Atmospheric Environment*, 177, 28–44. <https://doi.org/10.1016/j.atmosenv.2018.01.002>
- Li, Z., Guo, J., Ding, A., Liao, H., Liu, J., Sun, Y., Wang, T., Xue, H., Zhang, H., & Zhu, B. (2017). Aerosol and boundary-layer interactions and impact on air quality. *National Science Review*, 4(6), 810–833. <https://doi.org/10.1093/nsr/nwx117>
- Lim, S. N. (2015). Essentialising the convenient baba-nyonyas of the heritage city of Melaka (malaysia), 153–177. <https://doi.org/10.1057/9781137498601>
- Ma, S., Shao, M., Zhang, Y., Dai, Q., & Xie, M. (2021). Sensitivity of PM_{2.5} and O₃ pollution episodes to meteorological factors over the north China plain. *The Science of the total environment*, 792, 148474. <https://doi.org/10.1016/j.scitotenv.2021.148474>
- Majid, Jafari, A. J., Gholami, M., Fanaei, F., Arfaeina, H., et al. (2020). Association between meteorological parameter and PM_{2.5} concentration in Karaj, Iran. *International Journal of Environmental Health Engineering*, 9(1), 4.
- Montgomery, D., Peck, E. A., & Vining, G. (2021). *Introduction to linear regression analysis* (Sixth Edition). John Wiley & Sons, Inc.
- Ramli, N., Rubini, M., & Noor, N. M. (2024). Relationships between air pollutants and meteorological factors during Southwest and Northeast monsoon at urban areas in Peninsular Malaysia. IOP Conference Series Earth and Environmental Science, 1303(1), 012041–012041. <https://doi.org/10.1088/1755-1315/1303/1/012041>
- U. S. Environmental Protection Agency. (2023). Particulate matter (PM) basics. <https://www.epa.gov/pmpollution/particulate-matter-pm-basics>
- Wang, Y., Jia, C., Tao, J., Zhang, L., Liang, X., Ma, J., Gao, H., Huang, T., & Zhang, K. (2016). Chemical characterization and source apportionment of PM_{2.5} in a semiarid and petrochemical-industrialized city, northwest China. *The Science of the total environment*, 573, 1031–1040. <https://doi.org/10.1016/j.scitotenv.2016.08.179>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

COVID-19 Spread in Malaysia Using SIR Model

Nur Nadiah Az-zahraa Mohamed Sharfudeen^{1*}, Nurul Najihah Mohamad

^{1,2}Department of Computational & Theoretical Sciences (CTS), Kuliyyah of Science, International Islamic University Malaysia (IIUM), 25200 Kuantan, Pahang, Malaysia.

ARTICLE INFO

Article history:

Received 12 March 2025

Revised 14 May 2025

Accepted 18 May 2025

Published 1 September 2025

Keywords:

COVID-19

SIR Model

Markov Chain Model

Mathematical Modelling

Basic Reproduction Number

Transition Probability Matrix

DOI:

10.24191/jcrinn.v10i1.513

ABSTRACT

The COVID-19 pandemic had a significant impact globally. Negative impacts include the total number of losses in overall population size and economic decline. This study focuses on applying the simple Susceptible-Infected-Recovered (SIR) model to analyze COVID-19 cases in Malaysia for a time span of 100 days, from 1/5/2024 up to 8/8/2024. The key parts to gain the result can be divided into two which are data collection of daily COVID-19 cases in Malaysia from the website of Ministry of Health and solving the differential equations using R studio. From the SIR Model, the findings provide the estimation of transmission rate (β), recovery rate (γ), and a basic reproduction number (R_0), along with the graph of trends of COVID-19 in Malaysia for 100 days. From the values gained, this study aims to construct a Markov chain transition matrix to explain the disease spread more effectively.

1. INTRODUCTION

According to Achmad et al. (2021), COVID-19 is an emerging infectious disease, first found in late 2019, with the first patient being diagnosed on 31st December 2019 in Wuhan, China. Since then, it has been spreading rapidly across the country and eventually all over the world. The COVID-19 epidemic in Malaysia started with a small-sized outbreak of 22 cases in January 2020, mainly from imported cases. The first wave was quickly replaced by an even more significant outbreak, primarily secondary to local transmission, which led to 651 cases (Salim et al., 2020).

Salim et al. (2020) reported that there were three epidemic forecasting models used to make predictions for COVID-19 cases in Malaysia. The three models are the curve-fitting model with a probability density function and skewness effect, the stochastic model, and a system dynamic model. The stochastic process, which is modeled using Markov Chain, predicted a peak on May 20 and May 31, 2020, with 630,000 to 800,000 infections if the Movement Control Order (MCO) measure persists.

The pandemic has shown the necessity for continuous investigation and intervention to prevent the spread of the virus and the management of long-term effects. Hence, this study intends to see the spreading of COVID-19 in Malaysia by modeling it using a simple compartmental SIR model, where the parameters for the model are estimated from the real data.

^{1*} Corresponding author. E-mail address: nadiahsharfudeen@gmail.com
<https://dx.doi.org/10.24191/jcrinn.v10i2.513>

2. LITERATURE REVIEW

2.1 Disease spread simulation

According to Viboud et al. (2021), disease spread simulation is the application of mathematical or computational models to simulate and study the spread of infectious diseases in a population. These simulations are used to forecast the evolution of a disease in time, to assess the effects of different public health interventions, and to inform decision-making processes to control outbreaks in a proper way.

A cornerstone in disease spread model is the disease transmission dynamics, the field of research on how diseases move from one individual to another individual. Based on the website of CDC (Centers for Disease Control and Prevention, 2021) in 2021, disease transmission dynamics includes a consideration of factors involved in contact frequency, disease transmission probability, and number of new cases produced by one infected person that may occur.

An important number in the disease transmission dynamic is the basic reproduction number (R_0). It is the number that provides us the measure of how many new infections one infected individual can transmit to another individual (WHO, 2020). For example, if $R_0 = 2$, the minimum transmission by an infected individual is two cases per transmissions.

2.2 Previous study

Among the contributions in this field in the context of the COVID-19 pandemic, the paper by F. M. Omar, M. A. Sohaly, and H. El-Metwally (2024) provides a closed-form expression to estimate the Mean First Passage Times (MFPTs). MFPTs are also important in characterizing the time course of basic reproduction number (R_0) that describes the mean time for an infected individual to give the virus to the susceptible compartment. The use of this formula provides a more accurate estimate of the time for transmission events to take place, therefore improving the accuracy of virus spread forecasts. On the other hand, the authors also propose a state reduction approach to reduce the calculation of MFPTs, and the modelling process is more efficient by choosing a small number of states or steps. This method not only simplifies the computation but also makes it possible to perform faster and more realistic simulations. Additionally, they investigate whether the system's dynamics change significantly when new stages are introduced, such as distinct stages of infection and recovery lead to different relative levels and explain how such level changes lead to larger or smaller overall rates of virus transmission.

Another study by Wang and Mustafa (2023) have introduced an innovative Markov Chain Model designed for studying infectious diseases. This model is interesting because of its simplicity and efficiency, it can be controlled only by four parameters, even though it runs in an infinite state space. This advancement makes the model a versatile tool for analyzing complex disease dynamics without the need for extensive computational resources. Its streamlined design allows its use in democratized and usable ways in the context of many infectious disease settings, such as pandemic, i.e., COVID-19.

Moreover, the authors have performed many simulations methodology to cope with heterogenous infection distribution in various areas. By applying this method to six regions with high transmission rates, their study effectively analyses how varying regional characteristics influence disease spread. Their numerical study, based on a COVID-19 application as a case study, shows that it is possible to replicate the transmission dynamics that are characteristic of related infectious diseases using the model. This ability illustrates the power of the model to account for the dynamics of real-world disease transmission and its potential to improve epidemic prediction and public health plans.

2.3 Relationship between Markov chain model and SIR model

In this study, the disease spread model based on the SIR Model is constructed. SIR (Susceptible-Infected-Recovered) is a mathematical model that explains the population-level patterns of infectious diseases (Zenian et al., 2022). In Markov Chain Model, the principle of transition probabilities is applied

to the stochastic version of SIR Model. Since transition probabilities are typically represented in a transition matrix, let's set up a transition probability matrix in the context of SIR Model. Firstly, a transition diagram of the three states in the SIR Model is formed:



Fig. 1. The transition diagram of the three states in SIR Model

In this transition diagram, both the birth rate and death rate are assumed to be zero. The phase of moving from susceptible to infected state is represented by βIS where β is the transmission rate, I is the number of individuals currently in the infected state and S is the number of individuals currently in susceptible state. The term SI represents total number of potential interactions between the individuals in susceptible state and infected state. If βIS increases, the number of new infections in unit times also increases (Morris and Bjørnstad (2020)).

Logically, when the transmission rate is high, the infection is spreading faster, even with limited interactions between individuals in susceptible state and recovered state. On the other hand, observing the transition from infected state to recovered state, the transition is represented as γI where γ is the recovery rate that is proportional to the number of infected individuals I . It is crucial to understand that when the recovery rate is high, the rate of change from infective state to recovered state is high or the number of infective states moving to recovered state is bigger. The spread of the disease, βIS is non-linear since it depends on both state susceptible and infected, unlike the recovery phase, γI that is linear since it depends only on the infective phase (Bjørnstad, 2022). The balance between both transmission rate and recovery rate determines the basic reproduction number of a disease, R_0 :

$$R_0 = \frac{\beta}{\gamma} \quad (1)$$

Second step to construct a transition probability matrix for the SIR Model is listing all the probabilities in event. The key is to calculate the probability of moving from each state to another state. Since there are three states of "Susceptible", "Infected", and "Recovered", the matrix is supposed to be a 3×3 matrix. In Markov Chain Model, P_{ij} is known as the probability of moving or transitioning from initial state i to final state j . Therefore, listing the 9 probabilities of changing from the three states in SIR Model:

- (i) P_{SS} : the probability of susceptible individuals remains susceptible.
- (ii) P_{SI} : the probability of susceptible individuals becomes infected.
- (iii) P_{SR} : the probability of susceptible individuals to recovered.
- (iv) P_{IS} : the probability of infected individual becomes susceptible.
- (v) P_{II} : the probability of infected individuals stays infectious.
- (vi) P_{IR} : the probability of infected individuals to recovered.
- (vii) P_{RS} : the probability of recovered individuals become susceptible.
- (viii) P_{RI} : the probability of recovered individuals become infectious.
- (ix) P_{RR} : the probability of recovered individuals remains recovered.

Arranging these 9 probabilities into a matrix P :

$$P = \begin{bmatrix} P_{SS} & P_{SI} & P_{SR} \\ P_{IS} & P_{II} & P_{IR} \\ P_{RS} & P_{RI} & P_{RR} \end{bmatrix}$$

3. METHODOLOGY

3.1 Research design

This study implement the use of the SIR Model, which is originally considered as a deterministic model. This study is also non-experimental since existing data is used to model and predict outcomes of COVID-19 spread in Malaysia. According to Tang et al. (2021), in the SIR Model, the rate of change of the three states of “Susceptible”, “Infected”, and “Recovered” are depending on the factors of time. Therefore, in this research, the independent variable is time(t) in days. Meanwhile, the dependent variables are $S(t)$, $I(t)$ and $R(t)$, which the number of susceptible, infected, and recovered individuals at time t .

3.2 Data collection

The real-world data of COVID-19 in Malaysia is gathered from the official website of Info Centre by Ministry of Health (MOH). The data provide daily case reports with columns on region, number of people infected per day, number of people recovered per day, and number of deaths per day. The data has a time span of 100 days, from 1/5/2024 until 8/8/2024. The data is arranged according to the three states in SIR Model with number of infected individuals and number of recovered individuals per day. According to D'Ebarre (2019), to determine the number of susceptible per day, the number of infected individuals and recovered individuals is subtracted from the total number of population (N).

3.3 Model building

Three different states in SIR Model:

$S(t)$: Number of susceptible individuals at time t ,

$I(t)$: Number of infected individuals at time t ,

$R(t)$: Number of recovered individuals at time t .

Let the number of populations be N , therefore:

$$N = S(t) + I(t) + R(t). \quad (2)$$

The rate of changes of each state can be written as three different equations:

$$\frac{dS}{dt} = - \frac{\beta S(t)I(t)}{N}, \quad (3)$$

$$\frac{dI}{dt} = \frac{\beta S(t)I(t)}{N} - \gamma I(t), \quad (4)$$

$$\frac{dR}{dt} = \gamma I(t), \quad (5)$$

where β is transmission rate and γ is recovery rate.

3.4 Transition probability matrix

The agent-based modelling of the SIR Model will be explored. Agent-based modelling can be defined as computational simulations that attempt to model the behaviour of the individual within the environment (Singh, 2024). Firstly, computing the fractions of the three populations in susceptible, infected, and recovered as:

$$s(t) = \frac{S(t)}{N} \quad (6)$$

$$i(t) = \frac{I(t)}{N} \quad (7)$$

$$r(t) = \frac{R(t)}{N} \quad (8)$$

Then,

$$s(t) + i(t) + r(t) = 1. \quad (9)$$

Therefore, the rate of changes of each three states would become:

$$\frac{ds}{dt} = -\beta si, \quad s(0) = \frac{S(0)}{N} \quad (10)$$

$$\frac{di}{dt} = \beta si - \gamma i, \quad i(0) = \frac{I(0)}{N} \quad (11)$$

$$\frac{dr}{dt} = \gamma i, \quad r(0) = \frac{R(0)}{N} \quad (12)$$

To implement an agent-based simulation, assume the non-linear term is linear. Therefore, the rate of changes can be written in matrix such as:

$$\frac{d}{dt} \begin{bmatrix} s \\ i \\ r \end{bmatrix} = A \begin{bmatrix} s \\ i \\ r \end{bmatrix} \quad (13)$$

$$A = \begin{bmatrix} -\beta i & 0 & 0 \\ \beta i & -\gamma & 0 \\ 0 & \gamma & 0 \end{bmatrix} \quad (14)$$

Matrix A is the rate of change for a continuous time Markov Chain. Meanwhile, the matrix P is a simplified version of the transition probability matrix for SIR Model or the discrete-time Markov Chain. To prove the relationship between rate matrix A and transition matrix P , an exponential relationship between the two matrices is built, since matrix A is continuous to matrix P :

$$P = e^{A \Delta t} \quad (15)$$

For small Δt , the exponential can be approximated by the first-order Taylor expansion:

$$e^{A \Delta t} = I + A \Delta t \quad (16)$$

Where I = identity matrix and $A \Delta t$ = scales of rate matrix A by the time step Δt . Solving the first-order Taylor expansion:

$$A \Delta t = \begin{bmatrix} -\beta i \Delta t & 0 & 0 \\ \beta i \Delta t & -\gamma \Delta t & 0 \\ 0 & \gamma \Delta t & 0 \end{bmatrix} \quad (17)$$

$$P = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} + \begin{bmatrix} -\beta i \Delta t & 0 & 0 \\ \beta i \Delta t & -\gamma \Delta t & 0 \\ 0 & \gamma \Delta t & 0 \end{bmatrix} \quad (18)$$

$$P = \begin{bmatrix} 1 - \beta i \Delta t & 0 & 0 \\ \beta i \Delta t & 1 - \gamma \Delta t & 0 \\ 0 & \gamma \Delta t & 1 \end{bmatrix} \quad (19)$$

Next, it is important to understand that in the matrix A , it is an off-diagonal entries, which means the matrix represents the flows into a state. In matrix P , each rows represents the probabilities of leaving the state. Therefore, the diagonal term will remain the same, that is $P_{SS} = 1 - \beta i \Delta t$, $P_{II} = 1 - \gamma \Delta t$, and $P_{RR} = 1$. On the other hand, the off-diagonal terms will switch position with its opposite diagonal in the matrix, $P_{IS} = \beta i \Delta t$ will move to P_{SI} position and $P_{RI} = \gamma \Delta t$ will move to the position of P_{IR} . Hence, the transition probability matrix P is now:

$$P = \begin{bmatrix} 1 - \beta i \Delta t & \beta i \Delta t & 0 \\ 0 & 1 - \gamma \Delta t & \gamma \Delta t \\ 0 & 0 & 1 \end{bmatrix} \quad (20)$$

4. RESULT AND DISCUSSION

4.1 Estimation of transmission rate(β) and recovery rate(γ)

For this section, to obtain the estimation of transmission rate(β) and recovery rate(γ), R-programming language has been used (see Appendix 1). The coding is attached as in the Appendix 1. The results show the value of the transmission rate, $\beta = 0.006827$, recovery rate, $\gamma = 0.005924$, and basic reproduction number, $R_0 = 1.15242$. The value of transmission rate and recovery rate are closed to each other, indicating that during the 100 days, the rate of individuals getting infected from COVID-19 in Malaysia is close to the rate of its recovery. In short, the spread of the virus is in stable state since the infection and recovery rate operate at a similar pace. The basic reproduction number, $R_0 = 1.15242$, which is slightly above 1, which it indicates that the COVID-19 disease is still spreading in Malaysia but a lower rate. However, intervention steps still need to be taken since one infected individual can infect the virus over one other person.

According to Zenian et al. (2022), on the research of 'the SIR Model for COVID-19 in Malaysia' from January 2020 until June 2021, the values of transmission rate and recovery rate obtained were $\beta = 0.0162$ and $\gamma = 0.0069$ with a basic reproduction number, $R_0 = 2.35$. This occur since the cases of COVID-19 in Malaysia during 2020 to 2021 were still high and there was still limited knowledge about the virus and still fewer prevention strategies announced by the government.

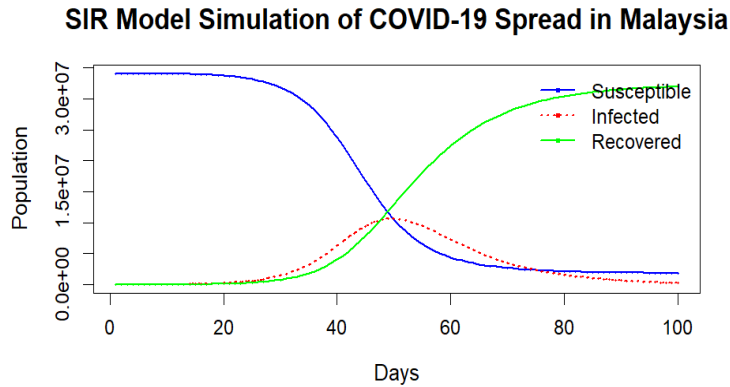


Fig. 2. The SIR model graph for $R_0 = 1.15242$

Fig. 2 illustrates the SIR Model on the COVID-19 pandemic in Malaysia from 1/5/2024 to 8/8/2024. The blue line represents the susceptible population of the 100 days period. Initially, the entire population is at the susceptible state except for those who are already infected before the first day, thus the curve begins at the maximum value of $N = 31\,400\,000$ minus the initial number of infected individuals. As the epidemic progress, the number of individuals in the susceptible state decreases gradually since some will move to the infected state. However, it will never reach zero since some individuals remain unexposed to the disease.

The red line represents the curve for the number of infected individuals, it is observed that the curve rises as the number of susceptible individuals decreases, indicating that the virus is spreading when both susceptible and infected individuals encounter each other. The number of infections proceed to increase until it reaches a peak value when it is about 50 days of the epidemic, this could happen due to the celebration of Eid al-Adha since the infection raises during the middle of June until early of July. After reaching the peak, the number of infected individuals starts to decline since more individuals are moving to recovered state. The green line, that is the curve for number of recovered individuals begins to increase slowly until it nearly approaches N .

4.2 Transition probability Matrix

For each time step, transition probability matrix P is:

$$P(t_{i+1}) = \begin{bmatrix} 1 - \beta \frac{I(t_i)}{N} & \beta \frac{I(t_i)}{N} & 0 \\ 0 & 1 - \gamma & \gamma \\ 0 & 0 & 1 \end{bmatrix}$$

$$P(t_{i+1}) = \begin{bmatrix} 1 - 0.006827 \frac{I(t_i)}{34\,100\,000} & 0.006827 \frac{I(t_i)}{34\,100\,000} & 0 \\ 0 & 0.994076 & 0.005924 \\ 0 & 0 & 1 \end{bmatrix}$$

Using the computed values of $I(t_i)$ from the simulation, the matrix P can be solved by finding the maximum value of the transition probability of moving from susceptible state to infected state, P_{SI} . The simulation result shows that the value of maximum transition probability from susceptible state to infected state, $P_{SI} = 0.006825$. Therefore, the final solution of matrix P is:

$$P(t_{i+1}) = \begin{bmatrix} 0.993175 & 0.006825 & 0 \\ 0 & 0.994076 & 0.005924 \\ 0 & 0 & 1 \end{bmatrix}$$

This matrix P provides the probabilities of the various transitions between states S , I , and R for each. The first row is transitions from the susceptible state, with the probability of remaining susceptible is 0.993175, which is remarkably high because most of the time an individual is not exposed to the infection due to the low fraction of infected individuals, . The probability of flow from susceptible to infected is 0.006825. The value is tiny since at this stage, the exposure to infected individuals is limited. The probability of moving from susceptible state to infected state is near to the transmission rate, $\beta = 0.006827$ since for small time steps, the probability of infection is proportional to the transmission rate. The likelihood of the infection is determined by the contact between susceptible and infected individuals, therefore, the higher the probability of transition from susceptible to infected state, the higher the transmission rate, β . The second row gives a probability of remaining infected that is moderately high at 0.994076 with most infected individuals remain in this state for more than one time step. The probability of transitioning to the recovered state is 0.005924, which directly reflects the recovery rate, $\gamma = 0.005924$. The third row is the transitions from the recovered state. The model assumes permanent immunity since the individuals that reach the recovered state remain there forever with probability 1, and there are no transitions back to either the susceptible or infected states as denoted by the zeros in the first and second column.

5. CONCLUSION

The findings indicate that the spread of COVID-19 was relatively kept under control during 2024, since lower rate of basic reproduction number implies lower transmission of the virus in the country. To minimize the basic reproduction number R_0 , it is a must to minimize the value of β and S while maximize the value of γ .

In this study, the Markov Chain Model is employed to estimate the transition probabilities between the epidemiological states of susceptible, infected, and recovered individuals. The estimated probability of transitioning from the susceptible to the infected state is 0.006825, indicating a relatively low likelihood of new infections. Conversely, the probability of transitioning from the infected to the recovered state is 0.005924, which, while still low, is comparatively higher than the infection rate. These findings suggest a downward trend in the number of active COVID-19 cases in Malaysia.

The findings suggest that in 2024, Malaysia had approached the stable disease control of COVID-19. This is due to the intervention strategy practices such as MCO, the use of my Sejahtera app, and massive vaccination program. Although progress has been made, continued awareness is necessary, as the virus still poses a risk of community transmission.

6. ACKNOWLEDGEMENTS/FUNDING

The authors would like to express their gratitude to the Department of Computation and Theoretical Sciences, Kulliyyah of Science, International Islamic University Malaysia (IIUM), Kuantan, Pahang for the support and facilities provided.

7. CONFLICT OF INTEREST STATEMENT

The authors agree that there is no conflict of interests regarding publication of this paper.

8. AUTHORS' CONTRIBUTIONS

Nur Nadiah Az-zahraa Mohamed Sharfudeen: Collected the data, prepared literature review, wrote the research methodology, made revision for result and finding. **Nurul Najihah Mohamad:** Conceptualisation and supervision, format analysis, writing-review and editing, validation.

9. REFERENCES

- Achmad, A. L. H., Mahrudinda, N., & Ruchjana, B. N. (2021). Stationary distribution Markov chain for Covid-19 pandemic. *Journal of Physics Conference Series*, 1722(1), 012084. <https://doi.org/10.1088/1742-6596/1722/1/012084>
- Bjørnstad, O. N. (2022). *Epidemics: Models and data using R*. Springer. <https://doi.org/10.1007/978-3-031-12056-5>
- Centers for Disease Control and Prevention (CDC). (2021). *Understanding the basics of epidemiology*. World Health Organization (WHO).
- D'ebarré, F., & Bonhoeffer, S. (2019). *SIR models of epidemics Level 1 module in Modelling course in population and evolutionary biology* (701-1418-00). Theoretical Biology Institute of Integrative Biology ETH Zürich.
- Morris, S. E., & Bjørnstad, S. O. N. (2020). *shinySIR: Interactive plotting for infectious disease models*. ResearchGate. https://www.researchgate.net/publication/349952415_shinySIR_Interactive_Plotting_for_Mathematical_Models_of_Infectious_Disease_Spread
- Omar, F. M., Sohaly, M. A., & El-Metwally, H. (2024). Susceptible-exposed-infectious model using Markov chains. *Journal of Nonlinear Mathematical Physics*, 31(1). <https://doi.org/10.1007/s44198-024-00167-3>
- Salim, N., Chan, W. H., Mansor, S., Bazin, N. E. N., Amaran, S., Faudzi, A. a. M., Zainal, A., Huspi, S. H., Hooi, E. K. J., & Shithil, S. M. (2020). COVID-19 epidemic in Malaysia: Impact of lockdown on infection dynamics. *medRxiv* (Cold Spring Harbor Laboratory). <https://doi.org/10.1101/2020.04.08.20057463>
- Singh, A. (2024, June 18). *Agent based modeling: Techniques and applications*. BuiltIn. <https://builtin.com/articles/agent-based-modeling>
- Tang, J., Bahnfleth, W., Bluysen, P., Buonanno, G., Jimenez, J., Kurnitski, J., Li, Y., Miller, S., Sekhar, C., Morawska, L., Marr, L., Melikov, A., Nazaroff, W., Nielsen, P., Tellier, R., Wargocki, P., & Dancer, S. (2021). Dismantling myths on the airborne transmission of severe acute respiratory syndrome coronavirus-2 (SARS-CoV-2). *Journal of Hospital Infection*, 110, 89–96. <https://doi.org/10.1016/j.jhin.2020.12.022>
- Viboud, C., Simonsen, L., & Flores, J. (2021). Disease surveillance and modelling in the time of COVID-19. *Journal of Infectious Diseases*, 224(1), 1-3. <https://doi.org/10.1093/infdis/jiab227>
- Wang, C., & Mustafa, S. (2023b). A data-driven Markov process for infectious disease transmission. *PLoS ONE*, 18(8), e0289897. <https://doi.org/10.1371/journal.pone.0289897>
- Zenian, S. (2022). The SIR model for COVID-19 in Malaysia. *Journal of Physics Conference Series*, <https://dx.doi.org/10.24191/jcrinn.v10i2.513>

2314(1), 012007. <https://doi.org/10.1088/1742-6596/2314/1/012007>

APPENDIX 1: ESTIMATION OF TRANSMISSION RATE(β) AND RECOVERY RATE(γ)

```

1 N<-34100000
2 Infected<-SIR[[3]]
3 Recovered<-SIR[[4]]
4 Susceptible<-SIR[[5]]
5 days<-1:100
6 head(SIR)
7
8 new_infections<-diff(Infected[days])
9 susceptible_current<-Susceptible[days[-length(days)]]
10 infected_current<-Infected[days[-length(days)]]
11
12 # Calculate beta
13 b<-(N*new_infections)/(susceptible_current*infected_current)
14 b<-b[is.finite(b) & b>0]
15 b_median<-median(b)
16
17 new_recoveries <- diff(Recovered[days])
18 infected_previous <- Infected[days[-length(days)]]
19
20 # Calculate gamma
21 gamma<-new_recoveries/infected_previous
22 gamma<-gamma[is.finite(gamma)& gamma>0]
23 gamma_median<-median(gamma)
24
25 #Calculate Basic Reproductive Number, R0
26 R0 <- b_median/gamma_median
27 cat("Median value of beta (transmission rate):", b_median, "\n")
28 cat("Median value of gamma (recovery rate):", gamma_median, "\n")
29 cat("Basic Reproductive Number (R0):" , R0)
30

```



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Youth Insights on Factors Influencing Housing Purchase and Rental Decisions

Mohamad Haizam Mohamed Saraf^{1*}, Muhammad Eidlan Hakimi Radzi²,
Syahmimi Ayuni Ramli³, Mohammad Fitry Md Wadzir⁴

^{1,2,3,4} Real Estate Management, Faculty of Built Environment, Universiti Teknologi MARA Perak Branch, Seri Iskandar Campus,
32610 Seri Iskandar, Perak, Malaysia.

ARTICLE INFO

Article history:

Received 19 March 2025

Revised 23 April 2025

Accepted 23 April 2025

Online first

Published 1 September 2025

Keywords:

Housing

Purchase Decisions

Rental Decisions

Mean Analysis

DOI:

10.24191/jcrinn.v10i2.518

ABSTRACT

Deciding on purchasing or renting a house are two significant financial decisions as it involves a large proportion of one's expenses, specifically to the youth. Although the property market price index in Malaysia statistically shows an increase every year as observed by the National Property Information Centre (NAPIC), demand for purchase still remains stagnant in the housing market. With inflation and the ringgit's depreciation, purchasing or renting a house possesses difficulties for youth to consider after they have been employed. Previous published literature on house decision is more focused on young professionals, investors and green homes, leaving a gap in the study of youth, where their perspectives are limited to assess. As it is prevalent to understand their perspectives, this research aims to fill a research gap by identifying the key factors that influence youth whether to purchase or rent in residential properties with. A quantitative approach was employed using a questionnaire for data collection. The results were obtained from mean analysis which indicated facilities and amenities and the country's economic stability attributes as highly influenced the decisions of house purchase or rental with 4.82 and 4.80 mean values respectively. The results contribute to the understanding of youth perspectives on housing and renting issues that might have an effect on the real estate market demand and supply policy and interventions, specifically on the residential real estate market signifying the present state of research.

1. INTRODUCTION

Youth in Malaysia face a tough housing market due to the constant challenges of purchasing a house. The National Youth Development Policy in Malaysia defines youth as individuals between the ages of 15 and 40, but Malaysia will enforce a new age limit for youth on January 1, 2026 (The Star, 2023). For youth, it

^{1*} Corresponding author. *E-mail address:* moham8841@uitm.edu.my
<https://dx.doi.org/10.24191/jcrinn.v10i2.518>

seems that the aspiration of owning a home appears unrealistic, as they think it's "significantly tougher" now compared to their parents' peak home-buying era. This is supported by a research conducted by the Institute for Youth Research Malaysia (IYRES) that observed the majority of youth (84%) do not own a house (Siti Shazwani et al., 2022). The study involved 4,116 respondents aged between 18-40 years. The struggles of youth in the housing market could be related to the decision that youths need to make whether they want to rent or own a home based on factors considered such as long and short term plans, living needs, household income and housing market conditions (Luong et al., 2023; Zheng et al., 2019). According to Hassan et al. (2022), the purchase decision on housing property is a conclusion after consideration of buying a house or real estate. With the current economic situation being less stable and the state of real estate market values increasing, it simultaneously heats up the real estate issue. Besides this, someone may need to consider many factors when deciding to buy a house because it is the most expensive expense for a household. This is because buying a house is a big decision that involves substantial financial commitment and long-term implications in life (Young, 2023). It is difficult for them to choose the right decision, and this is to ensure that new home owners do not have to bear their financial burden. Just as observed in the Malaysian House Price Index 2023 Report by the National Property Information Centre (NAPIC), the median house price in Malaysia for Q3 2023 is RM350,174 which is an increase from the median price of RM330,000 in 2022 (Valuation and Property Service Department, 2024). Likewise with the Housing Index in Q3 2023 which saw a small annual growth of 0.1% at 212.6 points (RM458,751 per unit). This is an increase from the Q3 2022 index point of 212.4. This shows that the increase in median property prices could be one of the factors in the perspectives of renting and owning houses in Malaysia. Because of that, renting is always assumed as an alternative to homeownership when the household lacks affordability to buy a house in the current market as the price skyrocketed. Several questions have arisen in this study, namely what are the factors that can influence the decision whether to rent or buy a house based on the current economic situation? The second question is what are the current youth perspectives to rent or purchase? Although some research has been done on the factors of renting or owning a house, the literature focused on green residential and employed respondents but is lacking on the youth's perspective. This research aims to fill the existing gap in the literature by specifically examining the perspectives of Malaysian youth regarding the decision to rent or buy a home in the context of rising house prices. By focusing on this demographic, the research will examine the unique factors influencing their housing decisions, which have been largely overlooked in previous research that primarily targeted employed individuals and green residential considerations. By analysing statistics, studies and the results of respondents' answers from questionnaire questions, this research aims to provide a comprehensive understanding of the dynamics surrounding the choice between renting or buying a home which in turn sets the stage for a detailed investigation of this critical decision in real estate today. The findings will also provide a nuanced perspective on the challenges and opportunities faced by young Malaysians in navigating the housing landscape especially in house ownership, thereby informing strategies that can better support their housing needs in the current housing market.

2. LITERATURE REVIEW

The range of attributes served as a reference for this study, which aimed to identify the attributes on house purchases and rental decisions.

2.1 Location suitability

Youth takes into account multiple aspects when assessing housing choices, such as location and available amenities or facilities. The location of a residence and its proximity to local amenities and points of interest are important elements that can affect decisions regarding home purchases (Mang et al., 2020). Location is viewed as the proximity of the home to public amenities or facilities that young civil servants find important (Muhammad Zamri et al., 2021). Examples of these facilities include hospitals, educational

institutions, public transportation, grocery stores, schools, and several others (Mustapa & Razib, 2023). The short distance is significant because young people do not require much time to access the amenities and services. In general, this indicates that they take into account different important aspects like accessibility, safety, environmental standards, community, and financial affordability when assessing housing choices. In addition, an ideal site for housing includes the availability of public transport and near to the workplace. There is no doubt that accessibility and transportation are very related to the selection of the location of the house that will be rented or purchased by the buyers. It is due to many things that need to be considered such as the distance from home to the place of work or study, cost of transportation and accessibility facilities such as public vehicles. Urban-rural environments are also important to purchase-rental decisions. Cities can be appealing to youth because of their lively job opportunities, social events, and availability of services. Although city living offers unparalleled accessibility, convenience, and employment opportunities, it comes with higher property prices, smaller living spaces, and increased noise and pollution levels (Fezili, 2023). Next is the social atmosphere of a neighbourhood, one's local involvement and one's social relations with neighbours as well as one's sense of community, privacy and safety are all important elements that provide social satisfaction for residents (Aksel & İmamoğlu, 2020). The emphasis on the location of areas with a sense of safety, neighbourhood environment and community involvement emphasise the importance of these elements in shaping their preferences for renting and buying a home.

2.2 Housing preferences

Housing design and concept plays an important attribute in purchase-rental decisions especially the development of the housing scheme area. If the area is deemed unsafe or inappropriate for living, youth may opt not to rent or buy there and instead select rental choices in better and safer locations. An instance can be seen in Taman Sri Muda, Shah Alam, which illustrates the housing development environmental problems that arose during the flash flood disaster. As development in Kota Kemuning evolved, new developments tend to utilise the current drainage system in Taman Sri Muda to channel rainwater into the river instead of constructing a new system which can potentially lead to disasters for residents or prospective purchasers or tenants. Youth prioritize making a strong first impression so they want their house design to suit their taste so that they can give off a good first impression to other people (Mustapa & Razib, 2023). They frequently seek properties that align with their lifestyle by providing energy efficiency, natural light, and proper ventilation. House orientation is an important matter in the formation of the youth in home ownership and renting. In Indonesia, researchers found that natural resources, Qiblah, thermal comfort, and disaster prevention influence the decision of purchasing or renting among potential buyers (Hasan et al., 2022). Youth now prioritise locations with reliable internet infrastructure, as it significantly impacts their capacity to work remotely, continue their education, and sustain social connections. Whether they intend to rent or purchase a home, it is essential to support their everyday activities. According to the 2022 property market report, up to 83% of Malaysian homebuyers preferred to have a high-speed internet infrastructure access and the Internet of Things (IoT) as important fixtures for their homes (Property Guru, 2023).

2.3 Financial stability

When determining whether to rent or buy, young people weigh their work security and possibility for future income growth as well as the current economic condition. Several factors have to be considered in terms of financial capability, including ability to pay the down-payment, securing a housing loan, repaying the loan with interest, and meeting other commitments throughout the loan repayment duration (Muhammad Zamri et al., 2021). When renting, people determine whether their monthly income is sufficient to cover rent, utilities, and other living expenditures. Economic stability is also important to foster a conducive environment for property transactions that affect the value of the property market. This is due to the fact that economic stability can influence the housing market and an individual's capacity to make financial choices. An unstable economic situation in a country can influence home ownership because of fluctuating interest rates and inconsistent property market prices. In Malaysia, the affordability is affected

by the unstable economy in terms of a weak currency, the depreciated value of the Malaysian Ringgit, and inflation (Liu & Ong, 2021). Even with the government initiatives aimed at offering affordable housing, it continues to be a challenge for the youth since currently their incomes fall short of covering mortgage payments. The motivation for home ownership is steadily decreasing, largely because of escalating prices which results in them having to rent a house as an alternative (Baskaran et al., 2020).

3. METHODOLOGY

The objective of this exploratory research, which needed a quantitative metric to determine and rank the factors influencing the decision to purchase or rent, is achieved by using a quantitative method. A questionnaire was designed, and the validity and reliability have been tested with 22 pilot respondents. The questionnaire has been amended according to the academic responses for validity. A Cronbach alpha (α) analysis was used for the reliability of the questionnaire items. Joy (2007) states that the minimum number of respondents for a pilot test with a sample size not more than 300 is between 15 to 30. Thus, a reliability test was established using a pilot test of 22 pilot respondents not included in the actual respondents' sample. The data gathered from the pilot test was run in SPSS Version 25. The result of reliability shows a Cronbach Alpha (α) value of 0.754, which is interpreted as good reliability (Joy, 2007; Tavakol and Dennick, 2011). Following this, the questionnaire was distributed to employed youth in Perak using simple random sampling. The estimated youth population in Perak is 10,574,000 (Department of Statistics Malaysia, 2024). The sample required 112 respondents, which was calculated using Slovin's formula with assumptions of a 95% level of confidence and a 10% margin of error. Finally, the collected data was analyzed for descriptive statistics using SPSS Version 25, encompassing mean analysis presented and ranked in the form of tables.

4. ANALYSIS

This research focuses on employed youth respondents. This suggests that respondents in this survey possess the maturity to comprehend the questions, ensuring accurate and trustworthy results. The data collected were assessed using frequency and mean analysis based on the responses of employed youth in Perak, Malaysia. Out of 112 questionnaires, 103 responses (92%) were valid and accepted for analysis, while 9 responses (8%) were deemed invalid because the respondents were either unemployed or exceeded the age limit for youth. Research by Sataloff and Vontela (2021) suggests that a response rate above 50% is considered favorable for social science research. In addition, according to Fincham (2008), a response rate of 60% or higher is considered acceptable in most research survey areas.

Table 1. Demographics

Item	Questions	Frequency	Percentage
Gender	Male	53	51.5%
	Female	50	48.5%
Education	Diploma	59	57.3%
	Bachelor's Degree	26	25.2%
	Master's Degree	18	17.5%
Occupation	Public	46	44.6%
	Private	17	16.5%
	Others	40	38.9%
Ownership status	Owned	37	35.9%
	Rent	66	64.1%
Salary	< RM4,849	78	75.7%
	RM4,850 - RM10,959	23	22.3%
	> RM10,960	2	1.9%

Source: Authors' Work (2025)

The gender distribution of the respondent is 51.5% male and 48.5% female. Majority of the respondents' education level is diploma with 57.3%, bachelor's degree with 25.2% and master's degree 17.5%. Next is the occupation of the respondents. 44.6% of the respondents are public servants, 16.5% of the respondents are private and others with 38.9% respondents. Following this is the ownership status of the current residence. 64.1% of the respondents are tenants and the remaining owned the house. Last but not least, the majority of the respondents with a salary of less than RM4,849 (75.7%), while 22.3% of the respondents with RM4,850 - RM10,959 and the remaining 1.9% with a salary of more than RM10,960.

In order to clarify the data acquired for Part B of the questionnaire, the three factors influencing youth to purchase or rent a house have been incorporated in a matrix table at literature review to confirm all the elements indicated are accurate and relevant to the questions of possible attributes that could influence the intention of youth in buying or renting a house. For mean values interpretations, mean scores can be categorised into five levels of interpretation to help understand trends and perceptions in the data (Hadiyanto, 2019). Mean values of 4.51 to 5.00 are "very high," indicating strong positive perceptions, while values between 3.51 and 4.50 are considered "high," reflecting favorable perceptions. Mean scores from 2.51 to 3.50 represent a "moderate" view, suggesting mixed perceptions, whereas scores of 1.51 to 2.50 are labeled "low," indicating negative perceptions. Finally, a score of 1.00 to 1.50 is categorised as "very low," signifying strong negative influences among respondents.

Table 2. Mean analysis

Factor	Attributes	N	Mean value	Standard Deviation	Interpretation
Location Suitability Average Mean Value: 4.456	To what extent do the available facilities and amenities influence your choice of rental or purchase property?	103	4.82	0.573	Very highly influence
	To what extent does the ease of access to workplaces influence your choice of rental or purchase property?	103	4.07	1.131	Highly influence
	To what extent does neighbourhood influence your choice of rental or purchase property?	103	4.62	0.794	Very highly influence
	To what extent does public transportation influence your choice of rental or purchase property?	103	4.66	0.635	Very highly influence
	When considering a house to rent or buy, to what extent does urban versus rural locations influence your choice?	103	4.11	0.569	Highly influence
Housing Preferences Average Mean Value: 4.313	Do housing developments like new construction and community amenities influence your decision to rent or buy a house?	103	3.43	1.193	Moderately influence
	Does the house design and concept influence your decision to rent or buy a house?	103	4.77	0.581	Very highly influence

	Is having high-speed internet a critical factor in your decision-making process for renting or buying a house?	103	4.74	0.626	Very highly influence
Financial Stability	How important is the financial commitment in influencing your decision to rent or buy a house?	103	4.21	1.160	Highly influence
Average Mean Value: 4.491	Does the economic stability of the country play a role in your decision to rent or buy a house?	103	4.80	0.566	Very highly influence
	To what extent does the house price influence your choice to rent or buy a house?	103	4.32	0.888	Highly influence
	How important are government initiatives on home ownership in influencing your decision to rent or buy a house?	103	4.63	0.592	Very highly influence

Source: Authors' Work (2025)

Factor 1 consists of locational suitability attributes. The attributes of facilities and amenities, neighborhood surroundings, and transportation networks are highly influential in housing rental and purchase decisions among youth, with mean values of 4.82, 4.62, and 4.66, respectively. The other two attributes that have a significant influence are workplace accessibility (4.07) and rural-urban locations (4.11). Factor 2 pertains to housing preferences. Findings show two attributes with very high mean values: house design (4.77) and concept, along with high-speed internet connection (4.74). The environment of the housing development is moderately influential, with a mean value of 3.43. The last factor is financial stability, as youth indicated that the country's economic stability (4.80) and government initiatives on home ownership (4.63) significantly influence their rental and purchase decisions. Financial commitment and housing prices also have a considerable impact, with mean values of 4.21 and 4.32, respectively. In the factor ranking, financial stability is ranked first with an average mean value of 4.49, followed by locational suitability with an average mean value of 4.46. The last factor, housing preferences, has an average mean value of 4.31.

5. DISCUSSION AND CONCLUSION

Based on the findings, facilities and amenities and country's economic stability indicated highly influenced attributes that affect youth decision on renting and purchasing a house. Challenges faced by youth especially the millennials and Gen Z in buying a house are tough compared to previous generations. For youth, the motivation of owning a house appears unrealistic as they think it is significantly more difficult now compared to their parents' prime years for purchasing houses. Importantly from the findings, the renting or buying decision depends on the price of the property and whether youth have enough money to buy a home in their preferred location (Ismail et al., 2023). Although youth could look into some of the government initiatives for housing, this blanket subsidy only benefits established working adults and young professionals. It is important that the government provides a more targeted approach to benefit the lower-income groups, especially the youth and fresh graduates. This thought is supported by a recent survey conducted by the Institute for Youth Research Malaysia that found a majority of young people could not afford to purchase a house and were only able to rent (Sohaimi, 2021).

This struggle with affordability is not just perception; hard data and statistics support it too. It is prevalent in the property market report that there is an increase in the median price of a house in Malaysia which is reported at RM320,000 in 2022 and increased to RM350,000 in 2023 (PropertyGuru, 2023). The increase in housing prices in recent years has exceeded the growth in wages, making cost-effectiveness a crucial concern. Despite that, there is a lack of affordable houses available in the marketplace. These intertwined factors generate obstacles that render homeownership more difficult to attain for this youth generation. Supported by a study conducted by Mahazril et al. (2023) and Ismail et al. (2021) which indicate that Malaysian youth often consider housing prices, facilities and amenities, environment, and housing preferences when making their housing decisions, according to the variables examined in their research.

The factor ranking findings also indicate that financial stability is the most important factor in housing decisions. This suggests that youth prioritise affordability when selecting where to live or to rent, underscoring the need for affordable housing initiatives. Meanwhile, the lower ranking of housing preferences implies that these are less critical than financial and location factors, prompting further exploration of how consumer preferences evolve with economic conditions and trends, such as a shift towards sustainable housing or youth are now prioritising experiences and flexibility over property ownership in housing. At the end of the day, the decisions should be comforting. If youth are not comfortable in taking mortgages for house purchase, rent is the next alternative ownership. By examining external factors such as interest rates, government housing initiatives and current market trends, it will provide a comprehensive understanding of the challenges faced by youth in driving the housing market.

6. LIMITATIONS AND RECOMMENDATIONS

These findings of the intention to purchase or rent a house attribute provide several contributions and implications, specifically to the research area related to the intention to purchase and rental of house. This study contributes to closing the literature gap on attributes of purchase-rental for youth. The groundwork for future research and informed decision-making in the field of purchase or rental decisions can be structured by bridging the knowledge gap. However, some limitations and recommendations could be highlighted. Although this study relies on quantitative data from distributed questionnaires, the incorporation of qualitative methods such as interviews with focused individuals can provide a broader view of their motivations and experiences regarding the choices they make whether to rent or buy a house. This approach can also provide a comprehensive picture of the dilemmas and situations of youth in facing these challenges. In addition, it is recommended to use a combination of both quantitative and qualitative research methodologies. In addition, the demographic scope must be expanded to diversify socio-economic backgrounds and geographical locations. The addition of this demographic scope can be beneficial in tracking changes in home ownership from the perspective of youth over time, especially in unstable economic conditions and policy changes. With the addition of this future study, it can contribute valuable insights and provide effective solutions based on the needs of youth demographics.

7. ACKNOWLEDGEMENTS/FUNDING

The authors thank the anonymous reviewers for their comments and suggestions that helped improve this manuscript.

8. CONFLICTS OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

9. AUTHORS' CONTRIBUTIONS

Mohamad Haizam Mohamed Saraf: Conceptualisation, methodology, formal analysis, validation, writing-original draft, writing-review and editing and supervision; **Muhammad Eidlan Hakimi Radzi:** Conceptualisation, investigation and writing-original draft; **Syahmimi Ayuni Ramli:** Data curation, writing-review and editing; **Mohammad Fitry Md Wadzir:** Visualisation and writing-review and editing.

10. REFERENCES

- Aksel, E., & İmamoğlu, Ç. (2020). Neighborhood location and its association with place attachment and residential satisfaction. *Open House International*, 45(3), 327–340. <https://doi.org/10.1108/ohi-05-2020-0035>
- Baskaran, S., Binu Kumar, S. K., Samtharam, S. R., Tangaraja, T., & Mahadi, N. (2020). home ownership motivation in Malaysia: A youth perspective. *International Journal of Academic Research in Accounting*, 10(2), 49–64. <https://doi.org/10.6007/IJARAFMS>
- Department of Statistics Malaysia. (2024, August 30). *Perak: Mid year population estimates*. <https://www.dosm.gov.my/portal-main/release-content/current-population-estimates-by-administrative-district->
- Fezili, F. (2023). *Malaysia property market report overview 2023*. Property Genie. <https://www.propertygenie.com.my/insider-guide/malaysia-property-market-report-overview-2023-ynRueSbNVuBqaKL3ToJ5jG>
- Fincham, J. E. (2008). Response rates and responsiveness for surveys, standards, and the journal. *American Journal of Pharmaceutical Education*, 72(2), 43. <https://doi.org/10.5688/aj720243>
- Hadiyanto, H. (2019). The EFL students' 21st century skill practices through e-learning activities. *Indonesian Research Journal in Education*, 461–473. <https://doi.org/10.22437/irje.v3i2.8036>
- Hasan M.I, Aminuddin. A. M. R., & Hazrina. H. B. M. (2022). Locality of building orientation in traditional Indonesian architecture: A systematic literature review. *Journal of Design and the Built Environment*, 22(3), 60–68. <https://ejournal.um.edu.my/index.php/jdbe/article/view/40725>
- Ismail, S., Manaf, A. A., Hussain, M. Y., Basrah, N., & Azian, F. U. M. (2021). Housing preferences: An analysis of Malaysian youths. *Planning Malaysia*, 19. <https://doi.org/10.21837/pm.v19i17.993>
- Joy, M. (2007). Research methods in education (6th Edition). *Bioscience Education*, 10(1), 1–2. <https://doi.org/10.3108/beej.10.r1>
- Liu, J., & Ong, H. Y. (2021). Can Malaysia's national affordable housing policy guarantee housing affordability of low-income households? *Sustainability*, 13(16), 8841. <https://doi.org/10.3390/su13168841>
- Luong, H. T., Tran, D. M., Nguyen, D. L. N., Nguyen, V. B., Le, A. T., & Pham, H. V. (2023). Determinants influencing housing-option decision of Gen Y: The case of Vietnam. *Journal of Distribution Science*, 21(7), 51–63. <https://doi.org/10.15722/JDS.21.07.202307.51>
- Mahazril, A., Nursyazwana, I., Kadapi, M., & Hussin, N. (2023). Housing preference among young people: A study from the perspective of university students in Malaysia. *Social and Management Research Journal*, 20(2), 41–51. <https://doi.org/10.24191/smrj.v20i2.24314>
- Mang, J., Zainal, R., & Mat Radzuan, I. S. (2020). Factors influencing home buyers' purchase decisions in <https://dx.doi.org/10.24191/jcrinn.v10i2.518>

- Klang Valley, Malaysia. *Malaysian Journal of Sustainable Environment*, 7(2), 77. <https://doi.org/10.24191/myse.v7i2.10265>
- Muhammad Zamri, N. E. M., Yaacob, M. ‘Aini, & Mohd Suki, N. (2021). Assessing housing preferences of young civil servants in Malaysia: do location, financial capability and neighbourhood really matter? *International Journal of Housing Markets and Analysis*, 15(3), 579–591. <https://doi.org/10.1108/ijhma-02-2021-0012>
- Mustapa, N. A., & Razib, M. R. (2023). Housing attributes that influence house buying decisions making by young working generations in Kuala Lumpur. *International Journal of Business and Technopreneurship (IJBT)*, 13(3), 223–230. <https://doi.org/10.58915/ijbt.v13i3.267>
- PropertyGuru. (2023, May 10). *Certificate Of Completion And Compliance (CCC) – What you need to know*. PropertyGuru Malaysia. <https://www.propertyguru.com.my/property-guides>
- Siti Shazwani, M. Y., Nur Diyana Nadirah, M. H., Siti Hanna Aisyah, N., & Ahmad Taufik, N. (2022). *Cabaran belia miliki kediaman sendiri*. Institute for Youth Research Malaysia. <https://iyres.gov.my/en/component/catalogue/item/432>
- Sohaimi, N. S. (2021). Should young professionals buy or rent a home? *International Journal of Academic Research in Business and Social Sciences*, 11(1). <https://doi.org/10.6007/ijarbss/v11-i1/9009>
- Tavakol, M., & Dennick, R. (2011). Making sense of Cronbach’s alpha. *International Journal of Medical Education*, 2, 53–55. <https://doi.org/10.5116/ijme.4dfb.8dfd>
- The Star Online (2023, 23 November). Youth age limit to be lowered to 30 by Jan 1, 2026. *The Star Online*. <https://www.thestar.com.my/news/nation/2023/11/23/youth-age-limit-to-be-lowered-to-30-by-jan-1-2026>
- Valuation and Property Service Department. (2024). *Property market report 2023*. Ministry of Finance, Malaysia. <https://laporan-pasaran-harta-tahunan/2023/>
- Young, B. (2023, December 21). *Challenges faced by millennials and Gen Z in buying a house and getting promoted*. EarnUp. <https://earnup.com/challenges-faced-by-millennials/>
- Zheng, S., Cheng, Y., & Ju, Y. (2019). Understanding the intention and behavior of renting houses among the young generation: Evidence from Jinan, China. *Sustainability*, 11(6), 1-18. <https://www.mdpi.com/2071-1050/11/6/1507>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Evaluating Students' Bread Preferences in UiTM Perlis: An Application of the Fuzzy Analytical Hierarchy Process

Khairu Azlan Abd Aziz¹, Wan Suhana Wan Daud^{2*}, Mohd Fazril Izhar Mohd Idris³, Nur Aqilah Kamal Azman⁴, Rizauddin Saian⁵

^{1,3,4,5}Faculty of Computer and Mathematical Sciences, Universiti Teknologi Mara, Perlis Branch, Arau Campus, 02600, Perlis.

²Institute of Engineering Mathematics, Universiti Malaysia Perlis, Kampus Pauh Putra, 02600, Arau, Perlis.

ARTICLE INFO

Article history:

Received 20 March 2025

Revised 18 May 2025

Accepted 30 April 2025

Published 1 September 2025

Keywords:

Fuzzy Analytical Hierarchy Process
Multi Criteria Decision Making
Bread Preference

DOI:

10.24191/jcrinn.v10i2.519

ABSTRACT

Nowadays, the food industry has expanded significantly, attracting many investors to invest in this business. Among these, the popularity of bread has been increasing, and many brands are available in the market. Bread has become a popular alternative food, especially among students, as it is convenient to store and carry. With numerous bread brands available, students have more options to choose from, give a challenge for bread manufacturers to compete with one another. This study aims to evaluate the factors influencing students' bread selection and determine the most preferred brand among UiTM Perlis students. It considers four criteria which are price, taste, packaging, and brand, while the alternatives bread brands are Gardenia, Massimo, Mighty White, and High Five. A multi criteria decision making method, the Fuzzy Analytical Hierarchy Process (FAHP), is applied to rank the preferred bread among UiTM Perlis students. The FAHP methodology involves data collection, consistency ratio measurement, and FAHP calculations. This approach is able to handle the subjective judgments of consumers. The results indicate that brand is the most influential factor, while the factor of price takes the lowest. Gardenia has become as the most preferred brand, followed by Mighty White, with High Five being the least favoured. This study provides valuable view of the students' bread preferences, helping industry players better understand consumer behaviour in the bread market while also improving their business and marketing strategies.

1. INTRODUCTION

Bread has been a preferred food in human civilization for centuries, serving as a primary source of food across various country. According to Skořepa and Pícha (2016), Mesopotamians and Egyptians eat bread for their daily meal. Traditionally, bread is made by mixing the grain flour with water. However, the industry of bread has evolved into variety of types and flavour to serve different consumer preferences worldwide. In Malaysia, the popularity of bread was started during the era of foreign trade and colonization.

^{2*} Corresponding author. E-mail address: wsuhana@unimap.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.519>

The British and Dutch as major colonial powers introduced the technique of making bread and the corresponding ingredients. Over the time, bread became an alternatives food of the Malaysian diet, influenced by the country's multicultural society. Chinese immigrants also played a role in promoting bread consumption, contributing to the emergence of various bread recipes that cater to different ethnic preferences. A multiracial and multicultural society in Malaysia has contributed to the diversification of bread production over the years. This development has provided consumers with a wider variety of choices.

The increasing demand for bread in Malaysia has been driven by industrial advancements after the independence. The expansion of commercial bakeries and large scale of bread production has made bread more available and affordable to the general public. With the growth of the bread industry, various brands have appeared and contribute to a huge number of breads. Started from traditional loaf bread, now the various types of bread with different flavours and shapes can easily found in the market. In the modern era, many artificial ingredients have been developed, which may affect the production of the bread industry. The high glycaemic index of bread can contribute to increased blood glucose levels, potentially leading to weight gain (Kourkouta et al., 2017). However, this occurs if consumers take bread in large quantities. The diverse selection of bread continues to attract consumers, particularly students who rely on it as a quick and convenient food option. As bread companies compete for market dominance, understanding consumer preferences becomes important. At UiTM Perlis for example, students are supplied with various bread brands through campus cafeterias and local markets. Vendors must ensure they stock preferred brands to minimize wastage and maximize sales since they are given the limited shelf life of bread. Hence, identifying the factors influencing students' bread selection and determining the most preferred brand can help vendors and manufacturers to improve their products.

In this study, Fuzzy Analytical Hierarchy Process (FAHP) is employed to identify the key criteria influencing bread selection and rank them to determine the preferred brand among students at UiTM Perlis. The study focuses on four criteria which are price, taste, packaging, and brand, followed with four major bread brands known as Gardenia, Massimo, Mighty White, and High Five. The data is collected from four student representatives, including Association and Club leaders, Student Representative Council members, and college committee members. The results of this study are anticipated to offer valuable insights for various stakeholders. As for example, the university's cafeterias can use the results to optimize their bread selection while bread manufacturers can enhance their products to meet the student preferences. Additionally, this research offers students an opportunity to understand and apply FAHP as a decision-making tool in evaluating consumer preferences. By implementing the FAHP method, this study may contribute to a better understanding of bread preferences among students, ensuring that vendors and manufacturers comply to consumer demands effectively while enhancing overall market strategies.

2. LITERATURE REVIEW

In this section, some theories behind the fuzzy numbers and FAHP is briefly introduced. Additionally, a review of previous studies on the topic and the hierarchical framework, including the criteria and alternatives considered in the study.

2.1 Fuzzy number

The idea of fuzzy set theory was introduced by Zadeh (1965). In this theory, values between zero and one are used to show how much something belongs to a set. Fuzzy sets are useful for describing uncertain or vague information in a meaningful way. From this idea, fuzzy numbers were developed to represent uncertain numerical values. One common type is the triangular fuzzy number, written as (l, m, u) , where l is the lowest possible value, m is the most likely value, and u is the highest possible value.

Definition 1: (Zadeh, 1965) A triangular fuzzy number (TFN) of $\tilde{A} = (l, m, u)$ has a membership function of $\mu_{\tilde{A}}$ provided by

$$\mu_{\tilde{A}}(x) = \begin{cases} \frac{x-l}{m-l}, & l \leq x \leq m, \\ \frac{u-x}{u-m}, & m \leq x \leq u, \\ 0, & \text{otherwise.} \end{cases} \quad (1)$$

where l and u stand for the fuzzy number's lower and upper bounds, respectively, while m is the median value. Fig. 1 illustrates the TFN in its standard form.

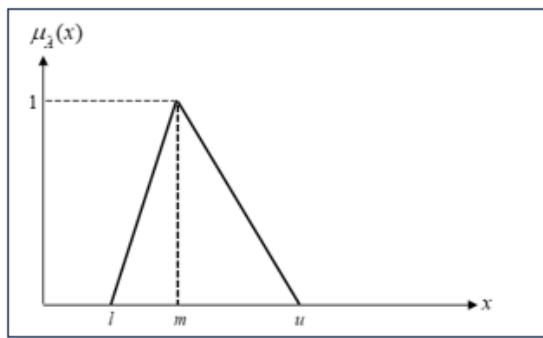


Fig. 1. Representation of a TFN (l, m, u)

2.2 Fuzzy analytical hierarchical process

In the 1970s, Thomas Saaty (Saaty, 1980) introduced the Analytic Hierarchy Process (AHP), a method that able to help in making decisions involving multiple criteria. AHP breaks down complex decisions into a hierarchy of goals, criteria, and options. Decision-makers then compare these elements in pairs to determine their importance, which helps rank the options. Until now, AHP has been widely used in various fields such as in business, healthcare, and government for solving many complicated problems. An extension of AHP that incorporates fuzzy logic is known as a Fuzzy Analytic Hierarchy Process (FAHP) which has been firstly proposed by Chang et al. in 1996 (Chang, 1996).

Unlike traditional AHP, FAHP uses a fuzzy triangular scale to handle uncertainties in subjective judgments. This adaptation is particularly useful in situations where decision-makers face ambiguity or imprecise data. In applying the FAHP method, data is collected through a questionnaire that facilitates pairwise comparisons of all boundaries and categories for analysis and ranking. This data is gathered from subject-matter experts who have the necessary qualifications to provide relevant insights, using a judgmental sampling technique. As noted by Saaty and Ozdemir (2015), FAHP primarily relies on expert opinions rather than a strict statistical approach. Therefore, the validity and consistency of judgments in AHP depend on the expertise of the participants in the specific field.

In FAHP, fuzzy number is often represented in reciprocal form of TFN, particularly when inverting a fuzzy pairwise comparison. Normally, the reciprocal fuzzy number is defined as

Definition 2: (Chandran, 2005) A triangular fuzzy number (TFN) of $\tilde{A} = (l, m, u)$, which $0 < l < m < u$.

The reciprocal of $\tilde{A}$ denoted as $\tilde{A}^{-1}$, is approximated by

$$\tilde{A}^{-1} = \left(\frac{1}{u}, \frac{1}{m}, \frac{1}{l} \right) \quad (2)$$

Over the years, the FAHP has been widely applied in various fields to solve complex decision-making problems involving multiple criteria. Halim et al. (2014) used FAHP to prioritize the essential criteria by ranking design attributes in Product Line Architecture (PLA). A year then, Brajkovic et al. (2015) used FAHP to evaluate the criteria weights based on students' online behaviours alongside with the Fuzzy TOPSIS. Besides that, FAHP has been employed in evaluating the in-flight service quality (Li et al., 2017) and also in selecting the passenger's aircraft types (Dožić et al., 2018). Additionally, Calabrese et al. (2019) utilized the FAHP to identify key sustainability issues, emphasizing the need for businesses to integrate sustainability into their strategies to mitigate environmental and societal impacts. Most recently, in 2023 and 2024, Idris et al. (2023) used FAHP in selecting the effective ways to prevent the COVID-19 spread, while Aziz et al. (2024) applied the FAHP in obtaining the ranking of factors to select the online shopping platform in Malaysia.

On the other hand, there are a few studies conducted on the criteria consumers consider when choosing a bread brand. In 2013, Nair (2013) highlighted numerous parameters for selecting bread. The study found that 'freshness' and 'quality' were ranked the highest, followed by 'taste,' 'softness,' and 'date of manufacturing,' while 'price' and 'appearance' ranked the lowest. Other criteria listed in this study are 'variety', 'availability' and 'brand name'. Furthermore, Gava et al. (2016) investigated additional criteria such as 'shelf life,' 'familiarity,' and 'weight control,' alongside the usual factors of price, shelf life, taste, freshness, familiarity, and weight control in evaluating the factors influencing bread choices. While, Eglite and Kunkulberga (2017) state in their study that the characteristic in choosing the bread is 'external appearances', 'producers', 'packaging design', 'expiry date', 'buying the same bread', and 'price'.

With various factors influencing consumer preferences in bread selection, it is also important to examine the most popular bread brands in Malaysia. Among the many established brands, the most well-known are Gardenia, Massimo, Mighty White, and High Five (Muhammad Fikri et al., 2022). As of now, no study has been conducted to evaluate the most popular bread brands in Malaysia using AHP. However, two studies have explored a somewhat similar scope using AHP. Soja and Melani (2021) examined the priority criteria and sub-criteria in selecting a bread marketing strategy at King's Bakery, while Shukriah et al. (2024) evaluated consumer perspectives on the marketing mix features of peanut bread products in Indonesia.

The following Fig. 2 illustrates the hierarchical framework used in this study which initiated based on the literature reviews. The hierarchical framework shows that, there are four key criteria of evaluation consists of price, taste, packaging, and brand that influence consumer choices. These criteria serve as the basis for evaluating the available options. Meanwhile, the second level presents the four bread brands considered which are Gardenia, Massimo, Mighty White and High Five. Each brand is assessed based on all four criteria, ensuring a comprehensive evaluation. The decision-making process relies on pairwise comparisons, where brands are evaluated against each other for each criterion. The FAHP method is used to systematically quantify subjective preferences, reducing bias and improving accuracy. This approach also helps rank the brands and determine the most preferred option based on consumer preferences.

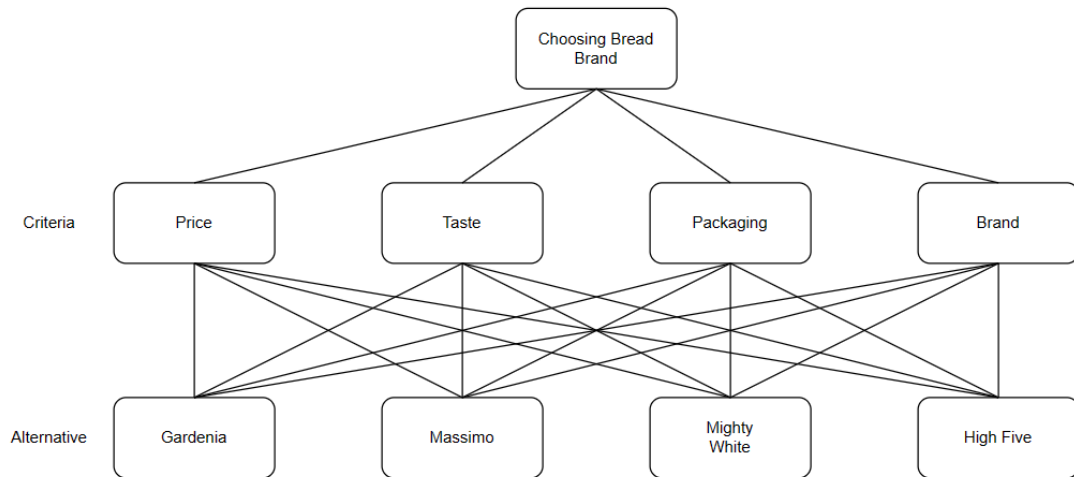


Fig. 2. Hierarchical decision model of the bread preferences

2.3 Criteria for the analysis

This section provides a brief overview of the criteria involved in this study.

2.3.1 Price

Price plays an important role in selecting a product to buy. Sometimes, the price reflects the quality of the product. Customers' perception of a product can often depend on its price. Albari and Indah (2020) discovered that the price of a product contributes to shaping its image.

2.3.2 Taste

Nordin and Teo (2024) identified taste as a key factor influencing adults' purchase intention for packaged food in Klang Valley. This factor is often related to the ingredients of the product. A good taste has a positive impact on the product and can also indicate its quality. The taste of packaged food may vary over time.

2.3.3 Packaging

Packaging is crucial when a company wants to sell its product. Good packaging represents the reputation and image of the company's product. Dutta and Sharma (2023) mentioned that quality and good packaging reflect the authenticity of the product. Moreover, it also helps employees promote the product.

2.3.4 Brand

The image of a brand can influence consumers to purchase a product. A strong brand is important, and manufacturers need to maintain its image to ensure its relevance for many years. Albari and Indah (2020) mentioned that manufacturers must consistently improve product quality and implement effective marketing strategies to sustain the brand.

3. METHODOLOGY

The methodology of this study is structured into three phases: data collection, consistency testing and weight calculation. Each phase comprises a series of steps, detailed in the subsequent sections.

3.1 Data collection method

The data used in this study consists of primary data which collected through a questionnaire. The questionnaire is constructed based on the previous study of FAHP (Harputlugil, 2018). In this study, four respondents consist of student committee and student clubs of UiTM Perlis are selected to complete the questionnaire. The respondents have some experience in organizing events for fellow students. Typically, the events they organized involve providing bread for breakfast for participants. Additionally, these students also are well-acquainted with the preferred bread brands among their peers.

3.2 Consistency test

Prior to calculating the weight for both the criteria and the alternative, it is necessary to conduct a consistency test to verify that the decision-maker's response is coherent. Inconsistency in the decision maker's answer examples include declaring A is more important than B while yet saying B is more significant than C. Despite this, the respondent afterwards said that C is more significant than A (Peng et al., 2021). This may be prevented if the consistency test is performed on all decision-makers. There are several steps in calculating the consistency ratio and the step is as follows:

Step 1: The first step is to make a pair-wise comparison matrix of the decision maker. The general form of pairwise comparison matrix is given as follows:

$$A = \begin{bmatrix} 1 & a_{12} & \cdots & a_{1n} \\ a_{21} & 1 & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \cdots & 1 \end{bmatrix} \quad (3)$$

where A is positive and symmetric matrix, since $a_{ji} = \frac{1}{a_{ij}}$ and $a_{ii} = 1$ for every $i, j = 1, 2, 3, \dots, n$. In other words, if the essential preferences a_{ij} is located in the upper triangle of the matrix, then the reciprocal value $a_{ji} = \frac{1}{a_{ij}}$ must be at the lower triangle or vice versa (Bozanic, et al., 2013).

Step 2: After the pair-wise comparison matrix is formed, the total for each criteria column is calculated, and the normalised matrix is found by dividing each cell by the total number of the respective column.

Step 3: The weight for each criterion is then determined by averaging the sum of each criterion.

Step 4: To find the weighted sum value, the first step is to multiply the original value of each cell in the judgement matrix by the normalised value that is found. Then total each of the criteria to get the weighted sum value.

Step 5: The ratio of the weighted criteria is found by dividing the weighted sum value by the weight criteria.

Step 6: After the ratio of the weighted criteria is found, the consistency index (CI) is calculated using Equation (4).

$$CI = \frac{\lambda_{\max} - n}{n - 1} \quad (4)$$

where $\lambda_{\max}$ is the largest eigenvalue of the comparison matrix, obtained by

$$\lambda_{\max} = \frac{\sum \bar{w}_i}{n} \quad (5)$$

which n is the number of alternatives.

Step 7: The consistency ratio can afterwards be calculated using Equation (6).

$$CR = \frac{\text{Consistency index (CI)}}{\text{Random consistency index (RI)}} \quad (6)$$

The greater the order of the matrix, the bigger the number in RI. The value in RI can be referred in the Table 1 according to the Kaganski et al. (2018).

Table 1. Random Consistency Index Matrix (RI)

Mean random consistency index										
n	1	2	3	4	5	6	7	8	9	10
RI	0	0	0.52	0.89	1.12	1.26	1.36	1.41	1.46	1.49

If the consistency ratio (CR) is less than 0.1, then the decision maker respond is consistent and acceptable (Afolayan et al., 2020). However, if the value of CR is greater than 0.1, it indicates a level of inconsistency that is too high, suggesting that the judgments may need to be revised or reconsidered to improve consistency (Saaty, 1980).

3.3 Weight calculation for criteria and alternatives

In this process, the ranking for both criteria and alternatives are determined using Buckley's approach (Ayhan, 2013). The steps are as follows:

Step 1: The criteria and alternative are compared by the decision maker using the linguistic term as shown in Table 1. The linguistic term is using the Saaty Scale (Saaty, 1980) and its corresponding fuzzy triangular number (Ayhan, 2013). The Saaty Scale is a fundamental tool used in the AHP and FAHP to compare the relative importance of different criteria or alternatives. This scale assigns numerical values ranging from 1 to 9, where 1 represents equal importance, while 9 signifies absolute importance of one criterion over another. Additionally, intermediate values (2, 4, 6, and 8) allow for better distinctions between levels of importance.

The linguistic term and the corresponding fuzzy triangular number are used in the construction of the pair wise comparison matrix. If the decision maker said that criteria 1 is fairly important compared to criteria 2, the fuzzy triangular scale are (4,5,6). On the other hand, the pair wise comparison will take the reciprocal of the triangular scale which are $(\frac{1}{6}, \frac{1}{5}, \frac{1}{4})$. The reciprocal values are used to maintain consistency in pairwise comparisons. This ensures logical coherence in the decision-making process.

In the FAHP approach, the Saaty Scale is extended using Triangular Fuzzy Numbers (TFNs) to handle uncertainty and imprecision in decision-making. Each scale value is represented as a fuzzy number in the form of (l, m, u) , where l (lower bound), m (middle value), and u (upper bound) define the range of possible judgments. For instance, a weakly important comparison (Saaty Scale = 3) is expressed as $(2,3,4)$ in fuzzy terms, reflecting a more flexible assessment rather than a fixed numerical value.

By incorporating Triangular Fuzzy Numbers into the Saaty Scale in Table 2, FAHP allows decision-makers to handle subjective judgments more effectively, capturing the uncertainty in human perceptions.

Table 2. Linguistic term and corresponding fuzzy triangular number

Saaty Scale	Linguistic Variable	Triangular Fuzzy Number	Reciprocal of the Triangular Fuzzy Number
1	Equally Important (Eq. Imp.)	(1,1,1)	(1,1,1)
3	Weakly Important (W. Imp.)	(2,3,4)	(1/4,1/3,1/2)
5	Fairly Important (F. Imp.)	(4,5,6)	(1/6,1/5,1/4)
7	Strongly Important (S. Imp.)	(6,7,8)	(1/8,1/7,1/6)
9	Absolutely Important (A. Imp.)	(9,9,9)	(1/9,1/9,1/9)
2	The Intermittent Value Between Two Adjacent Scales	(1,2,3)	(1/3,1/2,1)
4		(3,4,5)	(1/5,1/4,1/3)
6		(5,6,7)	(1/7,1/6,1/5)
8		(7,8,9)	(1/9,1/8,1/7)

The pair-wise comparison matrices that are obtained are then put in the comparison matrices according to the pair. The general comparison matrix is shown in Equation (7). The representation in each column indicates the k^{th} decision with the preference of the first criteria over the second criteria. The next element in the column is followed by the preference of the first criteria over the second criteria. The same can be said for the next row of the first column, which indicates the pair-wise comparison matrix of the preference of the second criteria over the first criteria.

$$\tilde{A}^k = \begin{bmatrix} \tilde{d}_{11}^k & \tilde{d}_{12}^k & \cdots & \tilde{d}_{1n}^k \\ \tilde{d}_{21}^k & \cdots & \cdots & \tilde{d}_{2m}^k \\ \cdots & \cdots & \cdots & \vdots \\ \tilde{d}_{n1}^k & \tilde{d}_{n2}^k & \cdots & \tilde{d}_{nm}^k \end{bmatrix} \quad (7)$$

Step 2: Since the decision maker preference is used for the calculation, if there exist more than one decision maker, the preference of the decision maker is averaged by:

$$\tilde{d}_{ij} = \frac{\sum_{k=1}^k \tilde{d}_{ij}^k}{k} \quad (8)$$

where k represents the number of experts and $\tilde{d}_{ij}^k = (l_{ij}^k, m_{ij}^k, u_{ij}^k)$. From the average value, the pairwise comparison is updated as shown in Equation (9).

$$\tilde{A} = \begin{bmatrix} \tilde{d}_{11} & \cdots & \tilde{d}_{1n} \\ \cdots & \ddots & \cdots \\ \tilde{d}_{n1} & \cdots & \tilde{d}_{nn} \end{bmatrix} \quad (9)$$

Step 3: The next step is to calculate the geometric mean of fuzzy comparison value for each of the criteria using the formula in the Equation (10).

$$\tilde{r}_i = \left(\tilde{d}_{i1} \times \tilde{d}_{i2} \times \cdots \times \tilde{d}_{in} \right)^{\frac{1}{n}} \quad (10)$$

where n is the number of factors. Subsequently, the vector summation of the geometric mean and its reciprocal are determined by means of the following formulas.

$$\begin{aligned} \sum_{i=1}^n \tilde{r}_i &= (\sum l_{\tilde{r}_i}, \sum m_{\tilde{r}_i}, \sum u_{\tilde{r}_i}) \\ \left(\sum_{i=1}^n \tilde{r}_i \right)^{-1} &= \left(\frac{1}{\sum u_{\tilde{r}_i}}, \frac{1}{\sum m_{\tilde{r}_i}}, \frac{1}{\sum l_{\tilde{r}_i}} \right) \end{aligned} \quad (11)$$

which $\sum_{i=1}^n \tilde{r}_i$ and $\left(\sum_{i=1}^n \tilde{r}_i \right)^{-1}$ represent the vector summation and reciprocal, respectively.

Step 4: The fuzzy weight of each criterion is calculated using Equation (12).

$$\begin{aligned} \tilde{W}_1 &= \tilde{r}_i \times (\tilde{r}_1 + \tilde{r}_2 + \cdots + \tilde{r}_n)^{-1} \\ \tilde{W}_1 &= lw_i + lm + uw_i \end{aligned} \quad (12)$$

This equation is calculated by first finding the vector summation of each criteria geometric mean. The reciprocal of the vector summation is then found before putting it in the Equation (12).

Step 5: The next step the defuzzification method. The number can be de-fuzzified by using the Centre of Area (COA) method using Equation (13).

$$COA = M_1 = \frac{lw_i + lm + uw_i}{3} \quad (13)$$

If $\sum M_1 > 0$, then it is normalized. Otherwise, it is not considered as a fuzzy number. Hence it needs to be de-fuzzified by using

$$N_1 = \frac{M_1}{\sum M_1} \quad (14)$$

The criteria with the highest weight are ranked first. As for the ranking of the alternative with respect to each criterion, an additional calculation step needs to be done. Equation (15) show the formula to obtain the ranking of the alternative with respect to each criterion.

$$R = \sum N_{criteria} \times N_{alternative-criteria} \quad (15)$$

4. RESULT AND DISCUSSION

The results are presented in the following order. First, the consistency ratio is checked for both the criteria and alternatives based on input from all respondents. Once the consistency ratio meets the required conditions, the weights are determined to establish the ranking of both criteria and alternatives.

4.1 Criteria consistency ratio

The consistency ratio for the criteria is tested for each expert by converting their questionnaire responses into a pair-wise comparison matrix. For instance, the pair-wise comparison matrix for Expert 1 is presented in this paper as follows:

Table 3. Expert 1 pair-wise comparison matrix

Criteria	Price	Taste	Packaging	Brand
Price	1	2	3	3
Taste	0.5	1	3	3
Packaging	0.33	0.33	1	3
Brand	0.33	0.33	0.33	1
Total column	2.1666667	3.6666667	7.333333333	10

Similar matrices were constructed for the other three experts. However, only the matrix for Expert 1 is presented as an example. After constructing the pair-wise comparison matrix, a normalized matrix for each expert is formed based on the normalized values for each criterion. This is achieved by dividing each criterion's value by the respective total value in the column for that criterion. The normalized matrix for the Expert 1 is shown in the following Table 4.

Table 4. Expert 1 normalized matrix

Criteria	Price	Taste	Packaging	Brand
Price	0.4615	0.5455	0.4091	0.3
Taste	0.2308	0.2727	0.4091	0.3
Packaging	0.1538	0.0909	0.1364	0.3
Brand	0.1538	0.0909	0.0455	0.1

After normalizing the matrix, the weight of each criterion is calculated by averaging the sum of each criterion. The weights for all criteria of Expert 1 are presented in the following Table 5.

Table 5. Expert 1 normalized matrix

Criteria	Weight
Price	0.429
Taste	0.3031
Packaging	0.1703
Brand	0.0976

Next, each value in the normalized matrix is multiplied by the previously computed weight of its associated criterion. The resulting values are then summed to obtain the weighted sum value.

Table 6. Expert 1 weighted sum value

Criteria	Price	Taste	Packaging	Brand	Weighted Sum Value
Price	0.429	0.6063	0.5108	0.2927	1.8388
Taste	0.2145	0.3031	0.5108	0.2927	1.3212
Packaging	0.143	0.101	0.1703	0.2927	0.7069
Brand	0.143	0.101	0.0568	0.0976	0.3984

Subsequently, the ratio of the weighted criteria is calculated based on the weighted sum value. The calculation is implemented by dividing the weighted sum value of each criterion with the respective weight of each criterion.

Table 7. Expert 1 ratio of weighted criteria

Criteria	Ratio of Weighted Criteria
Price	4.2861
Taste	4.3581
Packaging	4.1519
Brand	4.0836

Hence, the consistency index (CI) and consistency ratio (CR) for Expert 1 are obtained by using Equation (2) and (3).

Table 8. Expert 1 consistency index

Lambda Max	4.2199
Consistency Index	0.0733
Ratio Index (RI)	0.89
Consistency Ratio (CR)	0.0824

Based on the CR, the results are deemed consistent since the value is less than 0.1. The same processes are executed for all experts and presented in the following Table 9.

Table 9. Criteria consistency ratio

Expert	Consistency Ratio
Expert 1	0.0824
Expert 2	0.036
Expert 3	0.0723
Expert 4	0.0593

The CR is important for AHP and FAHP to make sure that the experts' answers make sense and acceptable. Based on the results shown in the Table 9, Expert 1 has a CR of 0.0824 (8.24%), Expert 2 has 0.0360 (3.60%), Expert 3 has 0.0723 (7.23%), and Expert 4 has 0.0593 (5.93%). Notably, Expert 2 has the highest level of consistency with a CR of only 3.60%, suggesting well-structured and reliable comparisons. Meanwhile, Expert 1 has the highest CR at 8.24%, but it still falls within the acceptable range. In overall, all the CR values are below 0.10 which indicates that the judgments made by these experts are consistent and valid for further analysis.

4.2 Alternative-criteria consistency ratio

In addition to the criteria pairwise comparison consistency test, consistency tests for alternatives with respect to each criterion pairwise comparison are also conducted. Hence, the second pairwise comparison is the pairwise comparison between each alternative and each criterion. The consistency tests performed are the consistency tests of alternative vs criterion pricing, alternative vs criteria test, alternative vs criteria packing, and alternative vs criteria variety. The procedures used to generate the consistency ratio for the criterion are used to compute the consistency ratio for each alternative criteria for each expert. The consistency test results are presented in the Table 10.

Table 10. Consistency ratio of alternative with respect to criteria

Expert	Expert 1	Expert 2	Expert 3	Expert 4
Consistency				
Alternative-Price	0.0646	0.0663	0.0925	0.0701
Alternative-Taste	0.0229	0.0232	0.0683	0.0937
Alternative-Packaging	0.0493	0.0319	0.0305	0.0174
Alternative-Brand	0.0692	0.0924	0.0649	0.0468

Table 10 shows the consistency exhibited by the experts for each alternative in relation to each criterion. When evaluating the consistency test of alternatives based on the price criteria, it is worth noting that Expert 1 achieved a consistency score of 0.0646, while Expert 2 obtained a slightly higher score of 0.0662. Expert 3 and Expert 4 both have consistency values of 0.0925 and 0.0701, respectively. When it comes to the consistency test for the alternative in relation to the taste criteria, all experts have consistency values ranging from 0.02300 to 0.0937, which is below the interval of 0.1.

The consistency tests for the alternative with respect to criteria packaging and the alternative with respect to criteria brand both show consistent results, falling within the ranges of 0.0174 to 0.0493 and 0.0468 to 0.0924, respectively. Based on the consistency ratio being less than 0.1, it can be inferred that the pairwise comparison of alternatives by all four experts, in relation to each criterion, is consistent.

Therefore, the next step involves aggregating expert judgments to derive the final weights for the criteria. This will be proceed using the geometric mean method, which helps combine individual expert opinions into a single, more representative set of priority weights. Once the aggregation is completed, the final ranking of criteria and alternatives can be determined, providing valuable insights for decision making.

4.3 Weight of criteria and alternative

At this phase, the calculation of FAHP is performed to obtain the weight for each criterion and alternative. The weight is determined by first calculating the fuzzy geometric mean for each respective calculation.

4.3.1 Weight of criteria

Determining the weight of criteria involves calculating the average contribution matrix. The process begins by converting the values in the pairwise comparison matrix of all experts to the fuzzy numbers based on the Saaty scale as shown in Table 2. The average contribution is then determined by calculating the mean of the fuzzy numbers provided by each expert. This information is summarized in the average contribution matrix, as shown in Table 11.

Table 11. Averaged contribution matrices for criteria

Criteria	Price	Taste	Criteria	Brand
Price	(1,1,1)	(0.14,0.45,0.54)	(1.61,1.88,2.16)	(0.62,0.90,1.20)
Taste	(2.75,3.50,4.25)	(1,1,1)	(2.63,3.17,3.75)	(2.31,2.82,3.34)
Packaging	(3.78,4.54,5.29)	(1.07,1.59,2.10)	(1,1,1)	(0.68,0.98,1.33)
Brand	(3.81,4.58,5.38)	(3.80,3.60,4.13)	(1.81,2.58,3.38)	(1,1,1)

Next, based on the averaged contribution matrices in the Table 11, the geometric mean for each criterion is calculated using the Equation (6) and presented as follows.

Table 12. Geometric mean of fuzzy comparison value

Criteria	RI		
Price	0.7986	0.9356	1.0892
Taste	2.0215	2.3651	2.7002
Packaging	1.2864	1.629	1.9629
Brand	2.1481	2.5548	2.9412
Total	6.2547	7.4845	8.6935
Reciprocal	0.1599	0.1336	0.115
Reciprocal increasing	0.115	0.1336	0.1599

Table 12 above presents the Relative Importance (RI) values for the four decision-making criteria: price, taste, packaging and brand. Each criterion is assigned a fuzzy RI value, represented by three bounds: lower (*l*), middle (*m*) and upper (*u*) or usually represented as (*l*, *m*, *u*) These values indicate the varying degrees of importance given to each criterion in the decision-making process. From the RI values, brand has the highest importance, with values of (2.1481, 2.5548, 2.9412) across the three fuzzy bounds. This confirms that brand perception is a dominant factor in purchasing decisions. Taste follows closely, with

values of (2.0215, 2.3651, 2.7002), showing that consumers also prioritize the sensory experience of the product. Packaging holds moderate importance, with RI values of (1.2864, 1.6290, 1.9629), while price has the lowest influence, with values of (0.7986, 0.9356, 1.0892). The total sum of the RI values for all criteria is (6.2547, 7.4845, 8.6935). To normalize these values, reciprocal values are calculated, which are (0.1599, 0.1336, 0.1150) for each respective fuzzy bound. These reciprocals are then used to adjust the RI values proportionally, ensuring that the final fuzzy weight distribution is consistent. Additionally, the reciprocal increasing values (0.1150, 0.1336, 0.1599) are applied in the defuzzification process to refine the ranking of the criteria.

From that, the relative fuzzy weight for each criterion is calculated by using the Equation (7), and presented as follows in Table 13.

Table 13. Relative fuzzy weight for criteria

Criteria	Weight		
Price	0.0919	0.125	0.1741
Taste	0.2325	0.3159	0.4317
Packaging	0.1479	0.2177	0.3138
Brand	0.2471	0.3413	0.4702

Next, the relative fuzzy weight for each criterion is defuzzified using the COA formula of Equation (8). As for the example, the COA for the criteria ‘price’ is shown, as follows.

$$COA(Gardenia) = \frac{(0.0919 + 0.1250 + 0.1741)}{3} = 0.1303$$

Table 14 presents the averaged defuzzified value for each criterion. Furthermore, the normalized value for each criterion needs to be determined since the total of the average defuzzified is more than 1. The calculation is using the Equation (9).

Table 14. Averaged and normalized fuzzy relative weight for criteria

Criteria	Average Defuzzified	Normalized
Price	0.1303	$0.1303/1.0365 = 0.1257$
Taste	0.3267	$0.3267/1.0365 = 0.3152$
Packaging	0.2265	$0.2265/1.0365 = 0.2185$
Brand	0.3529	$0.3529/1.0365 = 0.3405$
Total	1.0365	1

The total of the normalized values sums up to 1, which confirms that the weights for each criterion have been successfully obtained. The analysis of decision-making criteria using the FAHP method reveals that ‘brand’ holds the highest importance in consumer preference when selecting bread. With an average defuzzified weight of 0.3529 and a normalized weight of 0.3405, it is evident that consumers perceive brand reputation as a key factor influencing their purchasing decisions. This indicates that well-established brands have a strong impact on consumer trust and perception. Following closely, ‘taste’ is the second most important criterion, with a defuzzified weight of 0.3267 and a normalized weight of 0.3153. This suggests that consumers value the sensory experience of the product, emphasizing that a good taste enhances

satisfaction and loyalty. ‘Packaging’ ranks third, with a defuzzified weight of 0.2265 and a normalized weight of 0.2185. While not as critical as brand or taste, packaging still plays a significant role in attracting customers, preserving product quality, and influencing purchase decisions. Lastly, ‘price’ is the least influential factor, with a defuzzified weight of 0.1303 and a normalized weight of 0.1257. This finding suggests that consumers are willing to prioritize quality aspects such as brand reputation and taste over cost when selecting bread.

4.3.2 Weight of alternative-criteria

In addition to assessing the importance of the criteria, the evaluation also involves calculating the relative importance of each option with respect to each of the four criteria. The process for calculating these weights is identical to the procedure used to determine the weight of each criterion. The following Table 15 summarized the overall weight for each criterion.

Table 15. Weight of Alternative with Respect to Each Criteria

Alternative	Price	Taste	Packaging	Brand
Gardenia	0.3513	0.3913	0.2534	0.2718
Massimo	0.1471	0.2269	0.2267	0.1926
Mighty white	0.289	0.2343	0.3655	0.2781
High five	0.2125	0.1475	0.1543	0.2575

4.4 Ranking result

The procedure of obtaining the ranking of alternatives is performed by considering both the weights of the criteria and the aggregated results of the alternatives under each criterion. Table 16 presents the aggregated results for each alternative with respect to the criteria.

Table 16. Aggregated result for each alternative with respect to each criterion

Criteria	Weight	Scores of Alternatives with Respect to Criteria			
		Gardenia	Massimo	Mighty White	High Five
Price	0.1257	0.3513	0.1471	0.289	0.2125
Taste	0.3153	0.3913	0.2269	0.2269	0.1475
Packaging	0.2185	0.2535	0.2267	0.3655	0.1543
Brand	0.3405	0.2718	0.1926	0.2781	0.2575
	Total	0.3155	0.2051	0.2824	0.1946

The Table 16 presents the weighted scores of four bread brands Gardenia, Massimo, Mighty White and High Five based on the four key criteria which are price, taste, packaging and brand. Each criterion has a specific weight, indicating its relative importance in the decision-making process. As obtained previously, ‘brand’ has the highest weight, followed by ‘taste’, ‘packaging’ and ‘price’. Meanwhile, looking at the alternative scores which is the bread brands, Gardenia performs the best overall, with the highest scores in taste (0.3913) and price (0.3513), making it a strong choice for consumers. Mighty White stands out in packaging (0.3655) but falls behind in other categories. Massimo and High Five have lower overall scores, with High Five scoring the lowest in taste (0.1475) and packaging (0.1543), suggesting weaker consumer preference. Hence, the ranking results for both criteria and alternatives are determined based on their

respective weights, which obtained after utilizing the Equation (10). The example calculation is given for the Gardenia, which provides the highest ranking, as follows.

$$R(\text{Gardenia}) = (0.1257 \times 0.3513) + (0.3153 \times 0.3913) + (0.2185 \times 0.2535) + (0.3405 \times 0.2718) \\ = 0.3155$$

The overall results are presented in Table 17.

Table 17. Rank of each criteria and alternatives

Criteria	Scores	Rank	Alternative	Scores	Rank
Price	0.1255	4	Gardenia	0.3155	1
Taste	0.3152	2	Massimo	0.2051	3
Packaging	0.2185	3	Mighty White	0.2824	2
Brand	0.3405	1	High Five	0.1946	4

Given that ‘brand’ and ‘taste’ are the most influential criteria, the higher-ranked alternatives are likely perceived as stronger in these areas. Gardenia, holding the top position, likely benefits from a well-established and positive brand image combined with a favourable taste profile. Consumers likely associate Gardenia with quality and a satisfying sensory experience. It can also be said that even though ‘price’ has the lowest scores, Gardenia still excels in this aspect. Mighty White, ranked second, probably performs well in both ‘brand’ and ‘taste’, though perhaps not as strongly as Gardenia. It may have a slightly less prominent brand presence or a taste that is considered good but not exceptional compared to the market leader. It is also worth noting that since ‘packaging’ holds a low rank, these criteria would not be a factor for Mighty White. Massimo, ranking third, may face challenges in either brand perception or taste appeal. Consumers might be less familiar with the Massimo brand or find its taste less appealing than Gardenia or Mighty White. Alternatively, Massimo could be stronger in one area (e.g., taste) but weaker in the other (e.g., brand), resulting in a lower overall ranking. High Five, at the bottom of the ranking, likely struggles in both brand and taste. It may have a less recognizable brand identity and a taste profile that does not resonate as well with consumers. Also, it is worth noting that since both ‘price’ and ‘packaging’ holds a low rank, High Five excels at neither of these criteria. This clearly shows that in order for High Five to improve, the company should find a new branding that will allow them to compete with Gardenia and Mighty White. The ranking of alternatives appears to be heavily influenced by brand perception and taste appeal, with packaging and price playing secondary roles. To improve their position, lower-ranked alternatives should focus on strengthening their brand image and enhancing their taste profile to better align with consumer preferences. A more detailed investigation, perhaps involving consumer surveys or sensory evaluations, would be needed to confirm these hypotheses and identify specific areas for improvement.

5. CONCLUSION

Bread is an alternative food frequently consumed by students. With many brands available in the market, understanding the factors that influence their choices is important. This study applied the FAHP method to evaluate and rank these factors which involves subjective judgments. The result shows that the criteria brand is the most preferred factor for student to make a selection of bread, followed by taste, packaging, and price. A good brand not only produces a positive perception of the product but also convincing consumers of its quality and reliability.

Among the four bread brands observed Gardenia appeared as the most preferred alternative. This brand has built its own standard in the market and successfully winning the consumer trust. However, this result does not deny the quality of its competitors because each brand possesses its own strengths in many aspects. The findings show that bread manufacturers need to pay attention to improving their brand reputation. They must maintain customer trust and loyalty to their products. Investing in marketing and promotion can help consistently secure customer preference. There are several recommendations that can enhance the results of future research. Future studies could expand the scope by including additional factors such as texture, bread durability, ingredients, halal certification, and other elements that may influence consumer preferences. Furthermore, applying alternative models or techniques, such as fuzzy ELECTRE, could provide different perspectives on consumer choices. Additionally, future researchers may consider combining FAHP with other models, such as Fuzzy TOPSIS. This hybrid method could potentially enhance the results and make the findings more convincing.

By considering these improvements, future research can provide a better understanding of consumer preferences. The findings may benefit both consumers and manufacturers in making decisions regarding bread selection and product development.

6. ACKNOWLEDGEMENTS/FUNDING

The authors would like to acknowledge the support of Universiti Teknologi Mara (UiTM), Cawangan Perlis and Institute of Engineering Mathematics, UniMAP for providing the facilities support on this research. The authors would also like to express their gratitude to the anonymous referee for the constructive comments to improve this study.

7. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

8. AUTHORS' CONTRIBUTIONS

Khairu Azlan Abd Aziz: Supervision, conceptualisation, development of methodology, and formal analysis; **Wan Suhana Wan Daud:** Data interpretation, manuscript writing, and final editing.; **Mohd Fazril Izhar Mohd Idris:** Conceptualisation, methodology development, and formal analysis; **Nur Aqilah Kamal Azman:** Literature review, data collection, simulation of results, and validation; **Rizauddin Saian:** Language editing and contribution to the interpretation of results.

9. REFERENCES

- Afolayan, A. H., Ojokoh, B. A., & Adetunmbi, A. O. (2020). Performance analysis of fuzzy analytic hierarchy process multi-criteria decision support models for contractor selection. *Scientific African*, 9. <https://doi.org/10.1016/j.sciaf.2020.e00471>
- Albari, & Indah, S. (2020). The influence of product price on consumers' purchasing decisions. *Review of Integrative Business and Economics Research*, 7(2). 328-337.
- Ayhan, M. B. (2013). A fuzzy AHP approach for supplier selection problem: A case study in a gearmotor company. *International Journal of Managing Value and Supply Chains*, 4(3), 11–23. <https://doi.org/10.5121/ijmvsc.2013.4302>

- Aziz, K.A.A., Daud, W.S.W., Idris, M.F.I.M. & Zakaria, S. N. (2024). The ranking of factors in selecting the online shopping platform in Malaysia based on Fuzzy Analytic Hierarchy Process. *Journal of Computing Research & Innovation (JCRINN)*, 9(1), 31-42. <https://doi.org/10.24191/jcrinn.v9i1.393>
- Bozanic, D., Pamucar, D. (2013). Modification of the analytical hierarchical process method and its application in decision making in the defense system. *Technology*, 68(2), 327-334.
- Brajkovic, E., Sjekavica, T., & Volaric, T. (2015). Optimal wireless network selection following students' online habits using fuzzy AHP and TOPSIS methods. In *IWCMC 2015 - 11th International Wireless Communications and Mobile Computing Conference* (pp. 397-402). IEEE. <https://doi.org/10.1109/IWCMC.2015.7289116>
- Calabrese, A., Costa, R., Levialdi, N., & Menichini, T. (2019). Integrating sustainability into strategic decision-making: A fuzzy AHP method for the selection of relevant sustainability issues. *Technological Forecasting and Social Change*, 139(September 2018), 155-168. <https://doi.org/10.1016/j.techfore.2018.11.005>
- Chandran, B., Golden, B., & Wasil, E. (2005). A computational study of the analytic hierarchy process with exponential and triangular scales. *Computers & Operations Research*, 32(9), 2561-2574. <https://doi.org/10.1016/j.cor.2004.03.008>
- Chang, D. Y. (1996). Applications of the extent analysis method on fuzzy AHP. *European Journal of Operational Research*, 95(3), 649-655. [https://doi.org/10.1016/0377-2217\(95\)00300-2](https://doi.org/10.1016/0377-2217(95)00300-2)
- Dožić, S., Lutovac, T., & Kalić, M. (2018). Fuzzy AHP approach to passenger aircraft type selection. *Journal of Air Transport Management*, 68, 165-175. <https://doi.org/10.1016/j.jairtraman.2017.08.003>
- Dutta, D., & Sharma, N. (2023). Impact of product packaging on consumer buying behaviour: A review and research agenda. *Research Review International Journal of Multidisciplinary*, 8(7), 65-70. <https://doi.org/10.31305/rrijm.2023.v08.n07.009>
- Eglite, A., & Kunkulberga, D. (2017, April 27). Bread choice and consumption trends. *Foodbalt2017*. <https://doi.org/10.22616/foodbalt.2017.005>
- Gava, O., Bartolini, F., & Brunori, G. (2016). Factors in bread choice. *Italian Review of Agricultural Economics (REA)*, 71(1), 229-237. <https://doi.org/10.13128/REA-18642>
- Halim, S. A., Deris, S., & Zaki, M. Z. M. (2014). Fuzzy AHP based design decision for product line architecture. In *2014 8th Malaysian Software Engineering Conference* (pp. 119-124). IEEE. <https://doi.org/10.1109/MySec.2014.6986000>
- Harputulgil, T. (2018). Analytic Hierarchy Process (AHP) as an assessment approach for architectural design: Case study of architectural design studio. *Iconarp International Journal of Architecture and Planning*, 6(2), 217-245. <https://doi.org/10.15320/iconarp.2018.53>
- Idris, M.F.I.M., Aziz, K.A.A., Aziz, N. S. F. (2023). Selecting effective ways to prevent COVID-19 spread using the Fuzzy Analytic Hierarchy Process (FAHP) method. *Journal of Computing Research & Innovation (JCRINN)*, 8(2), 112-123. <https://doi.org/10.24191/jcrinn.v8i2.348>
- Kaganski, S., Majak, J., & Karjust, K. (2018). Fuzzy AHP as a tool for prioritization of key performance indicators. *Procedia CIRP*, 72, 1227-1232. <https://doi.org/10.1016/j.procir.2018.03.097>
- Kourkouta, L., Koukourikos, K., Iliadis, C., Ouzounakis, P., Monios, A., Tsaloglidou, & A. (2017). Bread and health. *Journal of Pharmacy and Pharmacology*, 5(11). <https://doi.org/10.17265/2328-2150/2017.11.005>

- Li, W., Yu, S., Pei, H., Zhao, C., & Tian, B. (2017). A hybrid approach based on fuzzy AHP and 2-tuple fuzzy linguistic method for evaluation in-flight service quality. *Journal of Air Transport Management*, 60, 49–64. <https://doi.org/10.1016/j.jairtraman.2017.01.006>
- Muhammad Fikri, M. A. A., Muhammad Amir, N. N. & Singh Sengar, S. (2022). Marketing strategies in peninsular Malaysia: Case study of Gardenia Barkeries Sdn. Bhd. *Asia Pasific Journal of Management and Education (APJME)*, 5(1), 94–107. <https://doi.org/10.32535/apjme.v5i1.1430>
- Nair, S. (2013). Assessing customer satisfaction and brand awareness of branded bread. *Journal of Business Management*, 12(2), 13–18. <https://doi.org/10.9790/487X-1221318>
- Nordin, Z. A., & Teo, S.S. (2024). Factors influencing consumers purchase intention of packaged food among adults in Klang Valley Malaysia. *Journal of Tourism, Hospitality and Culinary Arts*, 16(2), 131-149.
- Peng, G., Han, L., Liu, Z., Guo, Y., Yan, J., & Jia, X. (2021). An application of Fuzzy Analytic Hierarchy Process in risk evaluation model. *Frontiers in Psychology*, 12. <https://doi.org/10.3389/fpsyg.2021.715003>
- Saaty, T. L. (1980). *The analytic hierarchy process*. New York: McGraw-Hill.
- Saaty, T. L. & Ozdemir, M. S. (2015). How many judges should there be in a group? *Annals of Data Science, Springer*. <https://10.1007/s40745-014-0026-4>
- Skořepa, L., & Pícha, K. (2016). Factors of purchase of bread - prospect to regain the market Share? *Acta Universitatis Agriculturae et Silviculturae Mendelianae Brunensis*, 64(3), 1067–1072. <https://doi.org/10.11118/actaun201664031067>
- Soja, L. K., & Melani, A. (2021). Strategi pemasaran roti di King's Bakery dengan metode Analytical Hierarchy Process (AHP). *Industrika: Jurnal Ilmiah Teknik Industri*, 5(2), 81–85. <https://doi.org/10.37090/indstrk.v5i2>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Factory Simulation Construction Method and Implementation of Intelligent Manufacturing

Zheng Bin¹, Bai Yan^{2*}, Yang Soo Siang³, Tan Ming Keng⁴, Zhang Jing Song⁵

^{1,3,4,5}*Faculty of Engineering, University Malaysia Sabah, Kota Kinabalu, Sabah, 88400, Malaysia.*

^{2*}*Faculty of Mechanical Engineering, Beihua University, Jilin, 132013, China.*

ARTICLE INFO

Article history:

Received 28 March 2025

Revised 2 May 2025

Accepted 7 May 2025

Online first

Published 1 September 2025

Keywords:

Factory Simulation

Digital Twin

Intelligent Manufacturing

Simulation Modeling

Plant Simulation

DOI:

10.24191/jcrinn.v10i2.522

ABSTRACT

Construction methods and implementation of factory simulation of intelligent manufacturing discusses the key construction methods and realization ways of factory simulation in the field of intelligent manufacturing. As the core of the digital twin industry, it plays an important role in optimizing production and improving resource utilization. This paper expounds the construction method of factory simulation, including digital modeling, simulation parameter setting, advanced algorithm application and real-time data fusion. At the same time, it points out the challenges of model complexity and data accuracy in the simulation construction process and proposes corresponding solutions. This paper also takes Plant Simulation software as an example to analyze the implementation steps of factory simulation, such as model building, parameter configuration and process planning, emphasizing its value in optimizing factory layout, improving production process, and enhancing the scientific nature of decision-making. This paper comprehensively shows the key technical points and future development trend of factory simulation and provides theoretical support and practical guidance for the technological progress and industrial upgrading of intelligent manufacturing industry.

1. INTRODUCTION

With the rapid development of science and technology, intelligent manufacturing technology has become an indispensable part of industrial production (Tao et al., 2018). The use of automation, the Internet of things, artificial intelligence and other technologies can realize the production automation and intelligence, reduce human interference, and then improve the production efficiency and product quality (Li et al., 2017; Chander et al., 2022; Soori et al., 2023). As an important component of intelligent manufacturing, the construction method and implementation of digital twin factory involve a number of key fields, including digital modeling, virtual simulation, real-time monitoring, data analysis and optimization, cross-system integration and other (Liu et al., 2024; Segovia & Garcia-Alfaro, 2022). Through the digital modeling technology in this paper, the actual factory parts of the accurate digital processing, including production equipment, logistics system, personnel operation, etc., to establish a complete digital twin model. In the

^{2*} Corresponding author. *E-mail address:* baiyan@beihua.edu.cn
<https://doi.org/10.24191/jcrinn.v10i2.522>

digital twin model, the production process of the factory is simulated through the simulation technology, including the operation of the production equipment, product manufacturing and logistics and other links. At the same time, through the analysis and processing of the collected data, the digital twin model is optimized and adjusted to improve the production efficiency and product quality of the factory. This paper aims to improve the production efficiency of manufacturing industry, reduce operating costs, and optimize the allocation of resources. Therefore, the construction and development of intelligent digital twin factories promotes the transformation of industrial manufacturing to intelligent and digital, helps to enhance the core competitiveness of enterprises and accelerate the development process of industrial intelligence (Li et al., 2022; Wang & Luo, 2021).

2. INTELLIGENT PRODUCTION LINE OVERALL SCHEME DESIGN

2.1 Process design

This paper studies the firebrick detection line of an enterprise in Dalian, China. The firebrick detection line has realized the production operation, which covers the whole life cycle of firebrick from warehouse to finished products (Wahab, 2022). In order to meet the requirements of the product production process, the overall process includes the raw material delivery, transportation, mold release and cutting, steam maintenance, palletizing packaging and other links. Its complete operation process is shown in Fig. 1.

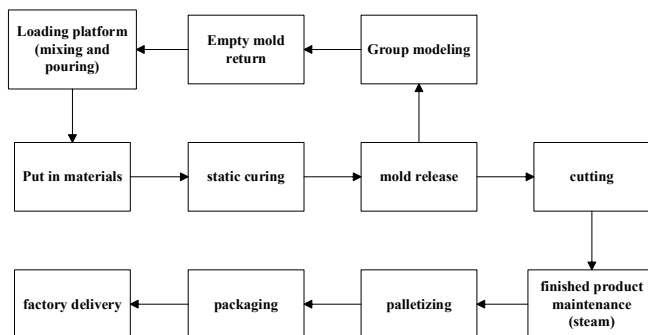


Fig. 1. Production line operation flow diagram

2.2 Intelligent production line planning scheme

The intelligent production line planning scheme is composed of layered design of equipment control layer, network layer, industrial cloud platform layer and application system layer (Qu et al., 2016; Wu et al., 2022). The equipment control layer is responsible for the overall planning, such as raw material control, valve monitoring, and data collection and remote control. The network layer is divided according to the hierarchical mode of the communication Internet industry, including the design of network, communication, cloud, and key basic capabilities (Peng et al., 2015). Network infrastructure mainly includes extranet, equipment, logistics tracking and energy monitoring. In terms of network form, it scientifically coordinates the planning of high-speed fiber network, wireless network and Iot sensing network; in the regional dimension level, the construction of high-speed fiber network and 4G / 4G + mobile communication network always follows the principle of comprehensive coverage of one network, and continuously promotes the commercial construction of 5G network (Stallings, 2015; Wen et al., 2021). The flow block diagram of the intelligent factory is shown in Fig. 2.

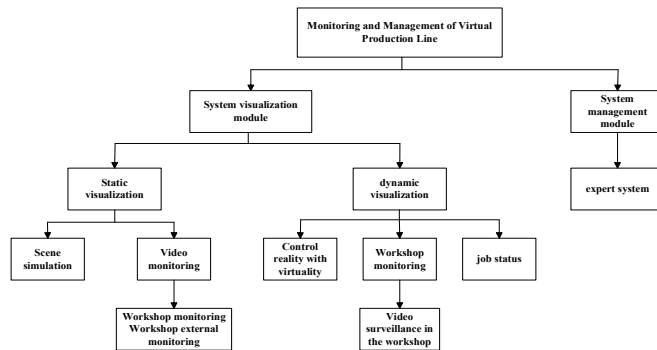


Fig. 2. Smart factory flow block diagram

3. HARDWARE STRUCTURE DESIGN OF INTELLIGENT PRODUCTION LINE

According to the required functions, the automatic detection line is designed for conveying system, central alignment system, separation system, 3D online measuring system, unloading system, marking and printing system, cardboard coating and installation system, and control system. As shown in Fig. 3. Overall design scheme diagram is the preliminary composition of firebrick automatic detection line. The delivery direction is shown in the arrow direction in Fig. 3.

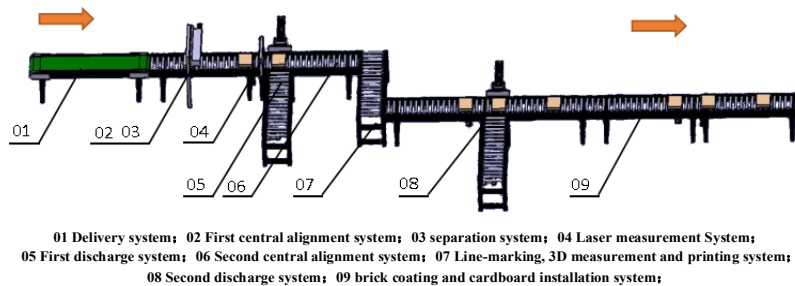


Fig. 3. Overall design scheme diagram

3.1 Conveying system

The conveying system runs through the whole brick detection line, which is the basis for all functional requirements of refractory bricks in the whole testing process. The conveying system mainly includes two types: the belt conveyor at the front end of the automatic detection line and the roller conveyor. As shown in Fig. 4 belt conveyor diagram and 5 roller conveyor diagrams are belt conveyor and roller conveyor respectively.

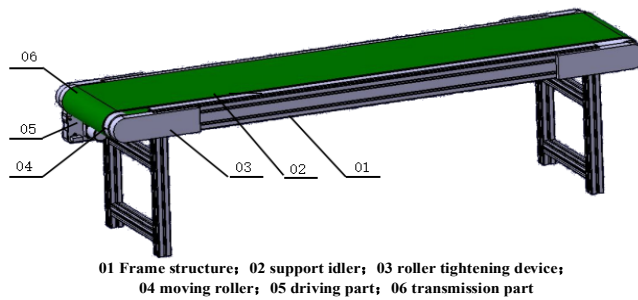


Fig. 4. Schematic diagram of the belt conveyor

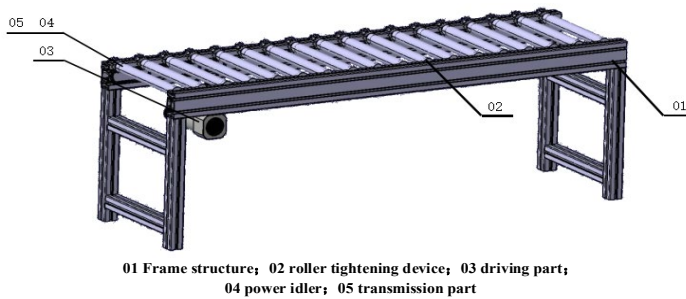


Fig. 5. Schematic diagram of the roller conveyor

3.2 Central alignment system, laser measurement system, separation system and unloading system

Fig. 6 central alignment system, laser measuring system, separation system and unloading system mechanical diagram, know the belt conveyor will brick into roller conveyor, into the central alignment system, by the system of the deviation of the brick, the mechanical alignment of the function of the motor gear, driven by the gear two relative rack movement, continue to transport forward into the separation system, the separation system is in a fixed position out baffle, let the brick stay, the function of the cylinder drive linear slide reciprocating movement (The diagram of the central alignment system structure and the separation system structure shown in Fig. 7) continue to move forward to the laser measurement system for data measurement and comparison. In the discharge system, the expansion cylinder will push out the unqualified refractory brick; the qualified one for alignment and separation will continue to be transported forward to the next station.

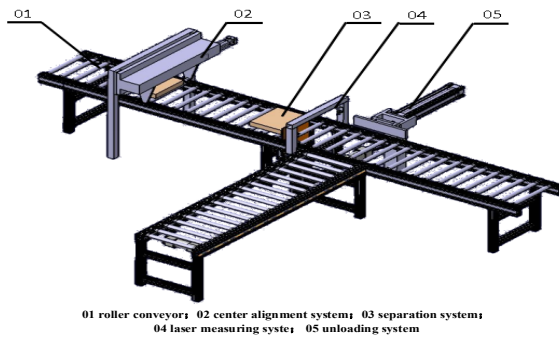


Fig. 6. Mechanical diagram of central alignment system, laser measuring system, separation system and unloading system

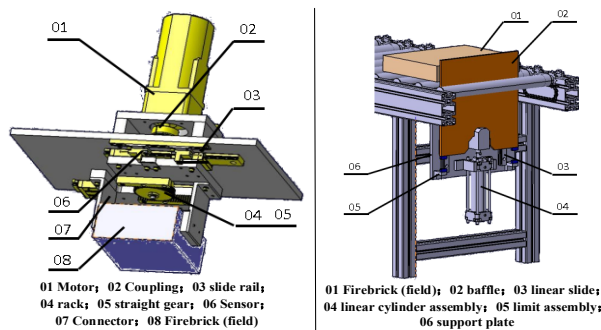


Fig. 7. Structure diagram of central alignment system and brief diagram of separation system

3.3 Line-marking, 3D online measurement and printing system

As shown in the schematic diagram of the marking, measurement and printing system in Fig. 8, the refractory brick qualified by laser measurement enters the second central alignment system, and the central alignment system pairs again. Then enter the marking, 3D online measurement and printing system.

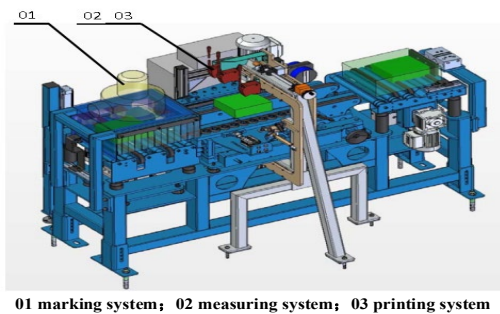


Fig. 8. Line-marking, measuring, and printing system

3.4 Brick bonding and cardboard mounting system

As shown in Fig. 9, the printed refractory brick is transported to the brick coating system by the roller conveyor, and the automatic coating equipment in the brick coating system is coated. The coating position and printing position are vertical end surface. the position relationship of brick printing and glue coating. The graphic arrow shows the direction of delivery. The finished bricks continue to the sticker system, attaching standard cardboard to the end of the bricks.

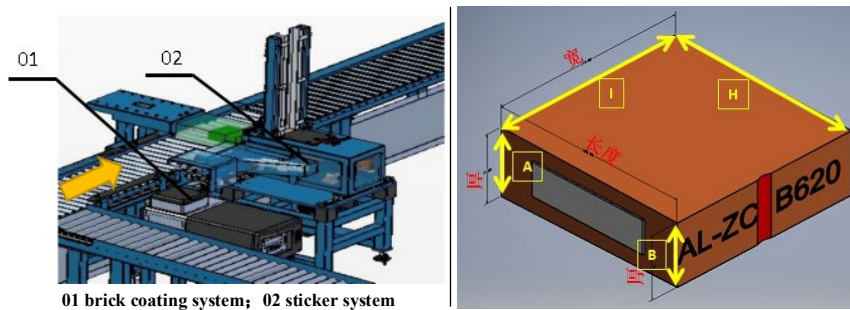


Fig. 9. Cardboard coating and installation and brick printing and coating position relationship

4. SOFTWARE PRODUCTION LINE SYSTEM DESIGN

This chapter studies two main parts: sorting and storage. First, the construction of the factory simulation environment based on Plant simulation software, then the design and construction of the raw material production unit and the raw material processing unit, and finally, the construction of the raw material storage unit.

4.1 Factory simulation environment construction based on plant simulation

4.1.1 The idea of the scene building

Before designing the production line model of the intelligent factory, the digital twin scene of the production line must be built according to the actual production process. The production line studied in this study needs to meet the sorting and transportation functions. Select the required model equipment from the digital twin model library and set up the production line in the Plant simulation 3 D modeling and simulation platform.

4.1.2 Scene building platform

Plant Simulation Is a software platform for the development of 3 D modeling and simulation display of digital twin systems, which can be used for the simulation and optimization of factories, production lines and production logistics processes. As part of the Siemens digital software Tecnomatix, it includes tools such as Process Designer, Process Simulation, and Plant Simulation. It is one of the key points of Siemens' digital manufacturing strategy. As a software that can simulate factories, production lines and logistics, Plant Simulation can quantify and verify all aspects of the production system, such as workshop layout, production logistics design and production capacity, and explore the optimization direction according to the simulation results, which can verify the effect of the implementation of the program before the implementation of the plan.

4.2 Simulation module design

4.2.1 Overall module design

The 2D / 3D layout diagram of the simulation model is as follows:

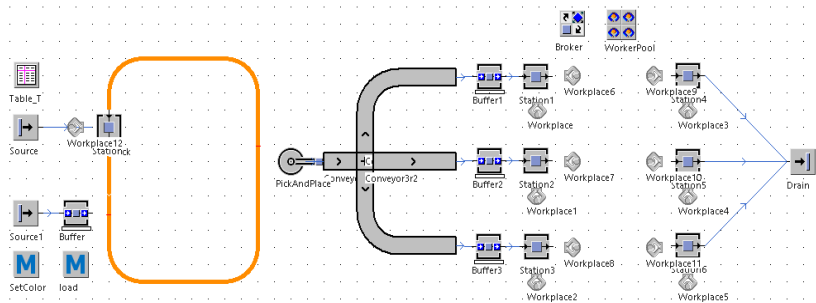


Fig. 10. Simulation model 2D layout diagram

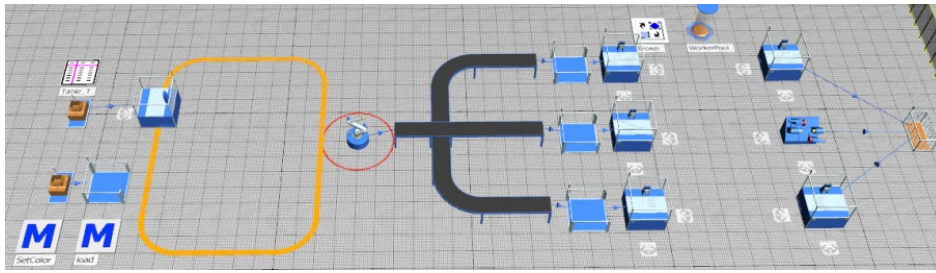


Fig. 11. Simulation model 3D layout diagram

In this model, the AGV trolley is generated from the Source on the left, and a buffer Buffer connects with the rail Track of the transport robot PickAndPlace, with two sensors on the rail Track. At the buffer zone Buffer, when the AGV car triggers the sensor, the car stops and loads the material, and the car continues to move when reaching the cargo capacity.

```
param SensorID: integer, Front: boolean, BookPos: boolean

@.stopped := true

while not @.full
    waituntil not Buffer.empty
    Buffer.cont.move(@)
end

@.stopped := false
```

Fig. 12. buffer Sensor controls

At the transport robot PickAndPlace, when the AGV car triggers the sensor, the car stops, the transport robot transfers the material to the production line from the AGV car, and the AGV car material is empty when the car continues to move.


```

param SensorID: integer, Front: boolean

@.stopped := true

var Destination:object := PickAndPlace

while not @.Empty
  var mu: object := @.cont
  if not mu.move(Destination)
    stopuntil mu.location /= @ and not Destination.isloading
  end
end

@.stopped := false

```

Fig. 13. PickAndPlace Sensor controls

4.2.2 Design of the transport and sorting area module

A transport robot PickAndPlace is placed on the right side of the AGV track, and a sensor is set at the AGV track. The function is to stop when the car triggers the sensor, and the transport robot transfers the material. When the material carried by the car is empty, the car continues to move along the track.

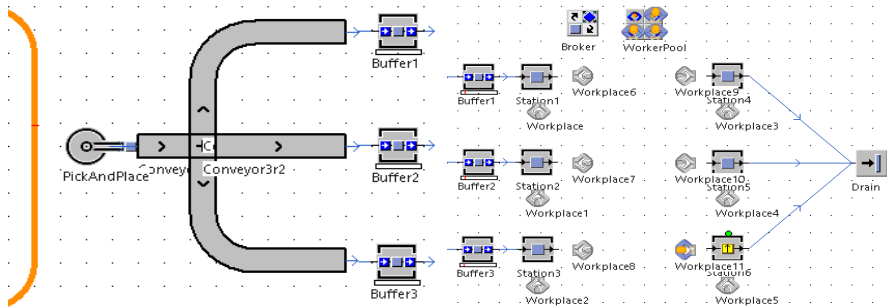


Fig. 14. Model layout

Fig. 15. Manufacturing zone model layout

A conveyor belt is placed next to the transport robot, and then a conveyor belt that can distinguish different materials and has a sorting function. The conveyor belts are connected to the other three ends.

4.2.3 Design of the raw material processing unit

Place a buffer behind each conveyor to prevent material from piling up on the belt. Two stations are placed after each buffer for material processing, with a distance between the two stations. A pool of workers is placed, in which ordinary workers and skilled workers with yellow circles are set up. The function of ordinary workers is to transfer materials from the first station to the next station.

One ordinary worker is set between each two stations, and a total of three ordinary workers are required in this model. Ordinary worker A1 transfers material from Station1 to Station4; ordinary worker A2 transfers material from Station2 to Station5, and ordinary worker A3 transfers material from Station3 to Station6. The technical workers with the yellow circle are responsible for the maintenance of the smart factory.

.Models.Model.WorkerPool.CreationTable						
*.Resources.Worker						
	Worker	Amount	Shift	Speed	Efficiency	Additional Services
1	*.Resources.Worker	1				A1
2	*.Resources.Worker	5				A2
3	*.Resources.Worker	1				A3
4	*.Resources.WORKER1	2				A0

Fig. 16. Worker pool creation table

Changing the value in the table above can change the number of workers and setting the Capacity (yield) value in WorkerSpace can carry multiple human jobs in a Workspace, that is, the number of Worker that can stay on the specified Workplace at any time.

4.3 Raw material storage unit design

The HBW (High Bay Warehouse) elevated cargo stereo warehouse in the Plant simulation toolbox is used to simulate the logistics and storage facilities in the actual factory. There are three components in this library object, namely WMS, roadway stacker and library location addressing control. These three components need to be used together to be effective. That is, each library component must be dragged to the modeling window to realize the operation and debugging of the three-dimensional warehouse.

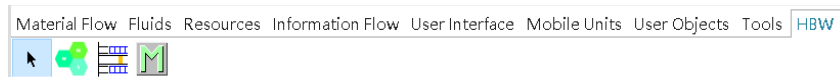


Fig. 17. Components in the HBW library

Under the interface of management, some very simple warehousing and warehousing strategies can be implemented, namely one by one, Random, Predefined racks and XYZ ranges. The key control information is the userSetTarget, which enables the scheduling of material storage. The important thing is a word, root.WMS.placeIntoStock, let WMS decide to put a material into the three-dimensional warehouse. This is the starting method of the stereo library, by which all other actions are triggered.

```
var foundFreePlace : boolean := root.WMS.placeIntoStock(@)

if not foundFreePlace
    self.openDialog
    messageBox("Could not find any free place in the racklanes", 1, 1)
    root.Eventcontroller.stop
end
```

Fig. 18. userSetTarget the algorithm

There are four warehousing and warehousing strategies in plant simulation software:

One by one: One place after another. Store the material from low to high, from left to right.

```

// Method Tasks:
// fill the rack lanes one by one
// Racklane is void if there was no free place found
//-----
param Pallet: object,
  byRef Racklane: object, byRef Side: string, byRef Column, Row: integer,
  product: string

var found: boolean := false
var finished: boolean := false

// just to make sure all racklanes have same utilization
var index: integer := lastIndex + 1
var count: integer := 1

if index > RackLanes.Dim then
  index := 1
end

```

Fig. 19. One by one source code

Random: That is, the randomly ground. Store the material in the stereo.

```

// Method Tasks:
// random search for a free place in the warehouse
//-----

param Pallet: object,
  byRef Racklane: object, byRef Side: string, byRef Column, Row: integer,
  product: string

var occupiedPlaces: integer

for var i := 1 to RackLanes.Dim
  Racklane := RackLanes[i]
  occupiedPlaces += Racklane.OccupancyLeft.sum
  occupiedPlaces += Racklane.OccupancyRight.sum
next

var totalPlaces: integer
for var i := 1 to RackLanes.Dim
  Racklane := RackLanes[i]
  totalPlaces += Racklane.NumberOfColumns*Racklane.NumberOfRows*2
next

var emptyPlaces: integer := totalPlaces - occupiedPlaces

```

Fig. 20. Random source code

Predefined: That is, the predefined. Store the material in an empty position close to the shortest path of the palletizing.

```

// Method Tasks:
// search for a free place in the predefined racks
//-----

param Pallet: object,
  byRef Racklane: object, byRef Side: string, byRef Column, Row: integer,
  product: string

var frame: object := current.~
var found: boolean := false

// get all predefined racks of the product
var PredefinedRacks: string[] := getPredefinedRack(Product)

if PredefinedRacks.dim>0 then
  // predefined rack found
  // check if there is a free place in one of the racks
  var index: integer := 1
  repeat
    Racklane:= frame.extendPath(PredefinedRacks[index])

    // get the first free place of the left rack
    var rack: object := Racklane.OccupancyLeft
    rack.setCursor(1,1)
    if rack.find(0) then
      var cLeft: integer := Rack.CursorX
      var rLeft: integer := Rack.CursorY
    end

    // get the first free place of the right rack
    rack := Racklane.OccupancyRight
    rack.setCursor(1,1)
    if rack.find(0) then
      var cRight: integer := Rack.CursorX
      var rRight: integer := Rack.CursorY
    end
  until found
end

```

Fig. 21. Predefined source code

XYZ ranges, namely the XYZ range. Store an item in the first free position in all the rack channels.

```

// Method Tasks:
// search for a free place in the warehouse
//-----
param Pallet: object,
byRef Racklane: object, byRef Side: string, byRef Column, Row: integer,
product: string

// get the range for the columns where the product will be stored
var Ranges: table := getProductRanges(product)
var frame: object := current.~
var found: boolean := false

if Ranges.YDim > 0 then
// in this table we will store the first free place of all rack lanes
var temp: table[object, string, integer, integer]
temp.create

for var i := 1 to Ranges.YDim loop
Racklane := frame.extendPath(Ranges[1, i])

var leftbound: integer := Ranges[2, i]
var rightbound: integer := Ranges[3, i]

// get the first free place of the left rack
var rack: object := Racklane.OccupancyLeft

rack.setCursor(1,1)
if rack.find({leftbound, 1}..{rightbound, *}, 0) then
temp.writerow(1, temp.YDim+1, Racklane, "left", Rack.CursorX, Rack.CursorY)
end

//get the first free place of the right rack
rack := Racklane.OccupancyRight

rack.setCursor(1,1)
if rack.find({leftbound, 1}..{rightbound, *}, 0) then
temp.writerow(1, temp.YDim+1, Racklane, "right", Rack.CursorX, Rack.CursorY)
end
next
end

```

Fig. 22. XYZ ranges source code

5. FACTORY SIMULATION OPERATION RESULTS

5.1 Overall factory simulation

The overall 3D diagram is shown in the figure below:

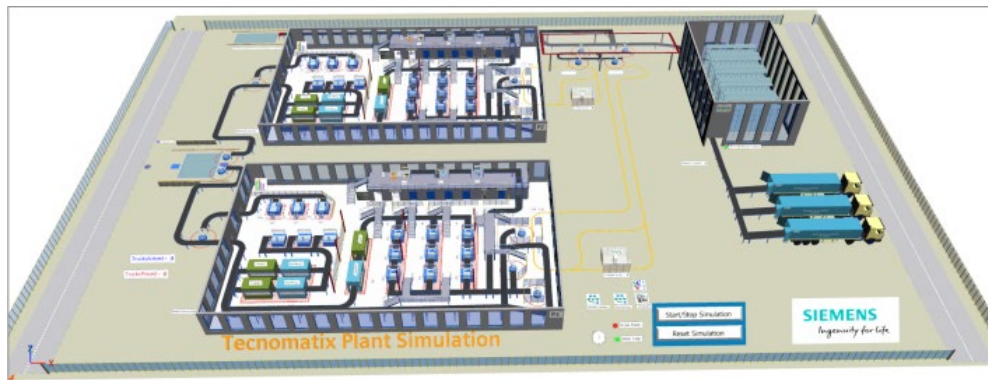


Fig. 23. Factory simulation overall 3D diagram

5.2 Raw material production unit

Fire brick raw materials through the conveyor belt transport to raw materials station for processing, processing by ordinary workers to the next station processing, processing through the conveyor belt successively through the spray machine and dryer spraying and drying, then by the conveyor belt transport successively successively through the paint machine and dryer for coating and drying.

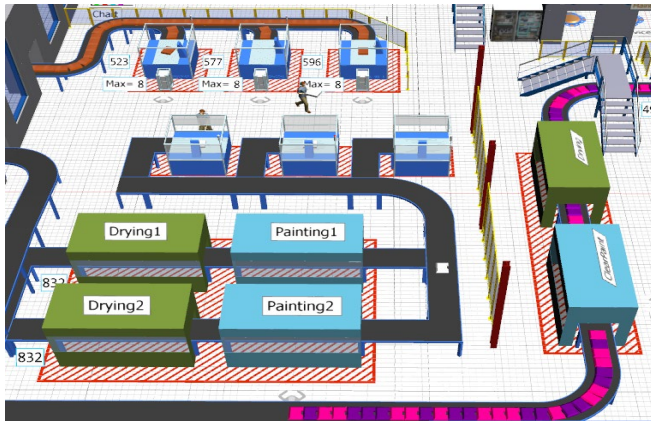


Fig. 24. Raw material production unit

5.3 Raw material processing unit

After the raw materials of refractory brick are sprayed with code and coating, they are transported by the conveyor belt to each station step by step. After the processing is completed, it will be transported to the transfer and sorting unit for packaging. Before the packaging, the workers will conduct manual sampling inspection on the processed raw materials to ensure the product quality.

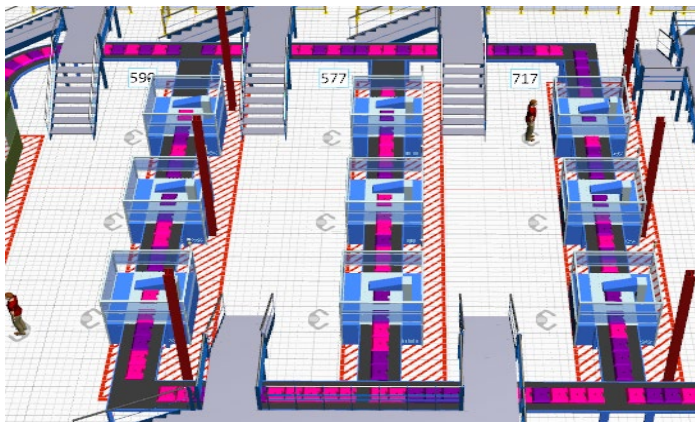


Fig. 25. Raw material processing unit



Fig. 26. Manual sampling inspection

5.4 Raw material storage unit

After packing the refractory bricks for stacking, they are transported to the three-dimensional warehouse for storage and transportation.

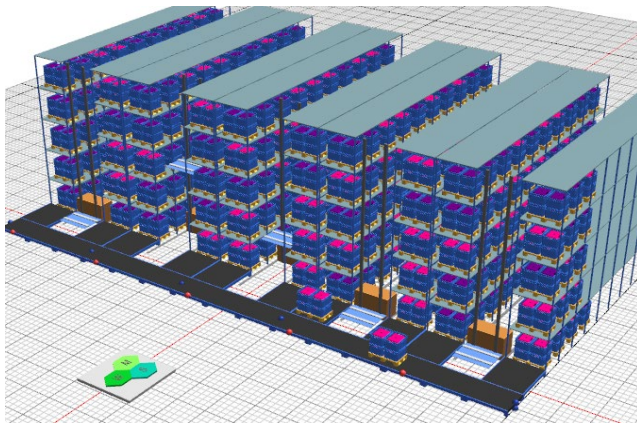


Fig. 27. Raw material storage unit

5.5 Computational results

Through the operation of the factory simulation, the conditions in the factory simulation, such as the processing situation, the station failure and maintenance situation of each station, are drawn into a diagram, and each part of the factory simulation is directly displayed.

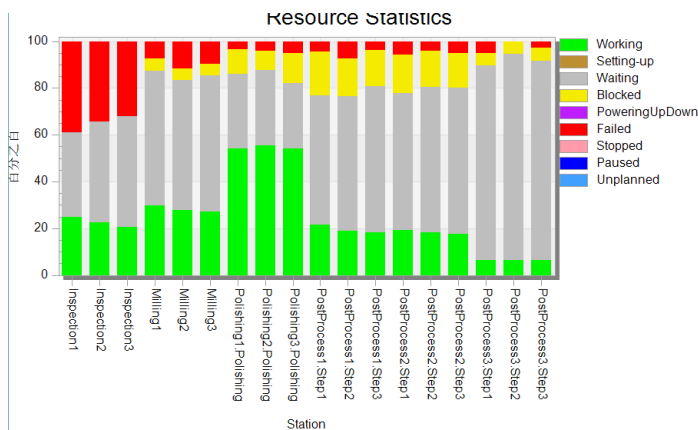


Fig. 28. Workplace diagram

6. CONCLUSIONS

Through simulation operation and debugging, the paper further optimizes the factory production line model, and the functions of material processing quantity, resource statistics and worker service statistics are added to make the simulation results closer to the actual operation of the factory. At the same time, this paper also verifies the advantages of three-dimensional warehouse in material storage and scheduling through the results of three-dimensional warehouse operation.

Despite the achievements made in this paper, there are still some improvements. First, the model should be further optimized to improve the accuracy and real-time performance of the simulation. Secondly, more simulation application scenarios should be explored, such as production line balance, equipment fault simulation, etc., to more comprehensively evaluate and optimize the production efficiency of the plant. In the future research, it is expected to further expand the application scope of simulation technology and

provide more support and help for industrial production. At the same time, we also hope to make more breakthroughs in the research and innovation of simulation technology and contribute more strength to the intelligence and automation of industrial production.

7. ACKNOWLEDGEMENTS/FUNDING

I would like to express my gratitude to the Faculty of Engineering of the University of Sabah in Malaysia and the Faculty of Mechanical Engineering of Beihua University in China for their support and assistance in this research. At the same time, I would also like to express my special thanks to all the participants and supporters throughout the research stage.

8. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

9. AUTHORS' CONTRIBUTIONS

Zheng Bin: System design, development and conduct testing; **Bai Yan** and **Yang Soo Siang:** Analysis of testing, supervision, paper review and editing; **Tan Min Ken** and **Zhang Jing Song:** Industrial expertise advisor.

10. REFERENCES

- Chander, B., Pal, S., De, D., & Buyya, R. (2022). Artificial intelligence-based internet of things for industry 5.0. In: Pal, S., De, D., Buyya, R. (Eds), *Artificial intelligence-based internet of things systems* (pp. 3-45). Springer. https://doi.org/10.1007/978-3-030-87059-1_1
- Li, B. H., Hou, B. C., Yu, W. T., Lu, X. B., & Yang, C. W. (2017). Applications of artificial intelligence in intelligent manufacturing: A review. *Frontiers of Information Technology & Electronic Engineering*, 18(1), 86-96. <https://doi.org/10.1631/FITEE.1601885>
- Li, L., Lei, B., & Mao, C. (2022). Digital twin in smart manufacturing. *Journal of Industrial Information Integration*, 26, 100289. <https://doi.org/10.1016/j.jii.2021.100289>
- Liu, Y., Feng, J., Lu, J., & Zhou, S. (2024). A review of digital twin capabilities, technologies, and applications based on the maturity model. *Advanced Engineering Informatics*, 62, 102592. <https://doi.org/10.1016/j.aei.2024.102592>
- Peng, M., Li, Y., Zhao, Z., & Wang, C. (2015). System architecture and key technologies for 5G heterogeneous cloud radio access networks. *IEEE network*, 29(2), 6-14. <https://doi.org/10.1109/MNET.2015.7064897>
- Qu, T., Lei, S. P., Wang, Z. Z., Nie, D. X., Chen, X., & Huang, G. Q. (2016). IoT-based real-time production logistics synchronization system under smart cloud manufacturing. *The International Journal of Advanced Manufacturing Technology*, 84, 147-164. <https://doi.org/10.1007/s00170-015-7220-1>
- Segovia, M., & Garcia-Alfaro, J. (2022). Design, modeling and implementation of digital twins. *Sensors*, 22(14), 5396. <https://doi.org/10.3390/s22145396>
- Soori, M., Arezoo, B., & Dastres, R. (2023). Internet of things for smart factories in industry 4.0, a review. <https://doi.org/10.24191/jcrinn.v10i2.522>

Internet of Things and Cyber-Physical Systems, 3, 192-204.
<https://doi.org/10.1016/j.iotcps.2023.04.006>

Stallings, W. (2015). Foundations of modern networking: SDN, NFV, QoE, IoT, and Cloud. *Addison-Wesley Professional*.

Tao, F., Qi, Q., Liu, A., & Kusiak, A. (2018). Data-driven smart manufacturing. *Journal of manufacturing systems*, 48, 157-169. <https://doi.org/10.1016/j.jmsy.2018.01.006>

Wang, P., & Luo, M. (2021). A digital twin-based big data virtual and real fusion learning reference framework supported by industrial internet towards smart manufacturing. *Journal of manufacturing systems*, 58, 16-32. <https://doi.org/10.1016/j.jmsy.2020.11.012>

Wahab, H. A. *Numerical study of brick using hybrid genetic algorithm*. chrome-extension://efaidnbmnnnibpcajpcglclefindmkaj/https://cdnx.uobabylon.edu.iq/undergrad_projs/D9MwvDloV0qWV1mXgRGyg.pdf

Wen, M., Li, Q., Kim, K. J., López-Pérez, D., Dobre, O. A., Poor, H. V., ... & Tsiftsis, T. A. (2021). Private 5G networks: Concepts, architectures, and research landscape. *IEEE Journal of Selected Topics in Signal Processing*, 16(1), 7-25. <https://doi.org/10.1109/JSTSP.2021.3137669>

Wu, Y., Dai, H. N., Wang, H., Xiong, Z., & Guo, S. (2022). A survey of intelligent network slicing management for industrial IoT: Integrated approaches for smart transportation, smart energy, and smart factory. *IEEE Communications Surveys & Tutorials*, 24(2), 1175-1211. <https://doi.org/10.1109/COMST.2022.3158270>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

PetalsPalette: Blossoming Knowledge by Bridging Technology and Nature through Augmented Reality Learning Application for Flower Identification

Nur Azmatun Farwizah Farid¹, Nurtihah Mohamed Noor^{2*}, Mohd. Fitri Yusoff³

^{1,2}*Faculty of Science Computer and Mathematical Sciences, Universiti Teknologi Mara Perlis Branch, Arau Campus, 02600 Arau, Malaysia.*

³*School of Creative Industry Management & Performing Arts, Universiti Utara Malaysia, 06010 Sintok, Kedah, Malaysia.*

ARTICLE INFO

Article history:

Received 11 April 2025

Revised 30 June 2025

Accepted 2 July 2025

Published 1 September 2025

Keywords:

Flower

Augmented Reality

Mobile Application

Learning

DOI:

10.24191/jcrinn.v10i2.524

ABSTRACT

This study focuses on learning different types of flowers using augmented reality (AR), which overlays digital content onto the physical environment. Studies regarding flowers or botanical species are complicated by a lack of adequate resources and tools, which puts barriers in place for educators and students who want to engage with the learning content. Current conventional methods of learning are not enough since some of the content may need further visualization of the real species or characteristics. Even though technological tools might be available, without necessary features and resources, particularly for visualization, the use will only be more distracting than helpful. Furthermore, providing an engaging user experience is crucial to ensure that learners effectively absorb the material. Our preliminary investigation has successfully validated the problem statement concerning lack of information and knowledge about flowers, which involves students in UiTM and educators. Thus, PetalsPalette, an AR application, was developed to provide opportunities for learners to learn about flowers through an AR which has been proven effective for learning. Five phases of methodology have been used, namely requirement analysis, design, implementation, testing, and maintenance. In the conducted testing, 33 respondents found that the PetalsPalette application is highly attractive, simple to understand, efficient, dependable, stimulating, and novel. PetalsPalette successfully enhanced flower identification for UiTM students through an AR-based mobile application. Future work aims to improve device compatibility and advance AR capabilities to create more inclusive and adaptable applications.

1. INTRODUCTION

Flowers have a profound ability to bring the enjoyment to those who are interested in them. This enjoyment often extends beyond the aesthetic beauty of the flowers that carry the hidden meaning and symbolism that

^{2*} Corresponding author. *E-mail address:* nurtihah@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.524>

can enhance the experience for enthusiasts (Kim et al., 2020). The study of flowering phenology and the anti-inflammatory properties of ornamental flowers adds to the understanding of the genuine interest individuals have in flowers (Łysiak, 2022). Moreover, the role of flower colour in angiosperm evolution draws attention to the evolutionary value of flowers and their visual characteristics (Narbona et al., 2021). These aspects collectively contribute to the widespread interest in flowers and the enjoyment derived from engaging with them.

However, learning about plants and flowers faces ongoing issues like the difficulty of identifying them, limited access to live examples, and challenges in seeing their features and traits (Łysiak, 2022; Chamidah et al., 2019; Narbona et al., 2021). Textbooks and still photograph are traditional ways of teaching that do not always show how the plants' or flowers' characteristics. This makes it challenging for learners to really grasp the appropriate information what they want to know and learn about the flower (Chamidah et al., 2019). Furthermore, since STEM (Science, Technology, Engineering and Mathematics) educations are critical and have been set to the highest consideration for a developing country to progress, such as Malaysia, it is crucial to maintain the interest to learn the knowledge. At the very least, to make individuals love the nature's flora and fauna more and save the environment. To maintain this, learners need more fun and engaging ways to learn. The lack of interactive and attractive tools might affect the interest of the individuals (Parsons et al., 2020), as they struggle to find information without any aid tool for the process. It also puts barriers in front of educators and students who want to engage with botany but lack the necessary tools and resources to effectively teach and learn about its species (Chamidah et al., 2019). The traditional methods, which involved manual techniques, are known to have limitations due to the lack of up-to-date and comprehensive information (Tonmoy & Islam, 2023).

Flowers are not only visually appealing but also possess intricate structures and features that have been the subject of extensive study. The morphological and anatomical study of flower structures has also been a subject of interest, shedding light on the diversity and characteristics of different flower morphs (Wulansari et al., 2022). Technology can be very important in improving understanding of and impressions of flowers. Studying and viewing flowers can be much more interesting and educational with the use of augmented reality (AR) technology. AR is a technology that integrates computer-generated content with the real world, allowing for the overlay of digital information onto the physical environment. This technology differs from virtual reality (VR) as it does not replace the real world but enhances it with digital elements. AR content can encompass various forms such as two-dimensional (2D) images, three-dimensional (3D) models, text, and real-time information, offering a novel and interactive experience (Vidak et al., 2022). AR has been found to have significant potential in various fields, including education, training, and industrial applications (Lam et al., 2021; Moesl et al., 2022; Wong et al., 2022).

The identification of flowers using AR involves the real-time recognition of flower species based on their physical and biological characteristics, integrating the real world with electronic information. This innovative approach offers significant advantages to researchers, particularly in the field of botany, by enabling the seamless combination of physical and digital realities for accurate flower identification (Setiyaningsih et al., 2021). The utilization of AR technology for flower identification has the potential to revolutionize botanical research and speed up species identification, offering a novel and efficient method for researchers to analyse and classify various flower species based on their unique characteristics (Marchenko & Kuzovkina, 2021; Rajesh Kumar & Subhashini, 2023; Setiyaningsih et al., 2021). A study by Wilujeng et al. (2019) looks closely at how AR technology can be used to create multimedia learning materials about plant structure, especially for the Wijaya Kusuma flower, aimed at college students. The findings demonstrated the successful creation of a prototype of the Wijaya Kusuma flower as a virtual object inserted into the real environment in real-time. However, the study also highlighted the need for further development to improve the application by integrating AR with mobile applications.

2. METHODOLOGY

Waterfall methodology has been used in this study. The waterfall methodology is a sequential software development process that consists of five distinct phases: requirements, design, implementation, testing, and maintenance, with each phase building upon the deliverables of the previous phase. The first phase, requirements, involves gathering and documenting the requirements. The second phase, design, focuses on creating the architectural and detailed design of the system based on the gathered requirements. The third phase, implementation, involves the actual coding and unit testing of the software. Finally, the maintenance phase involves making modifications, correcting defects, and enhancing the system based on user feedback and changing requirements. Fig. 1 shows the research activity diagram for waterfall methodology that involves five phases with activities, tools, or techniques and the deliverables.

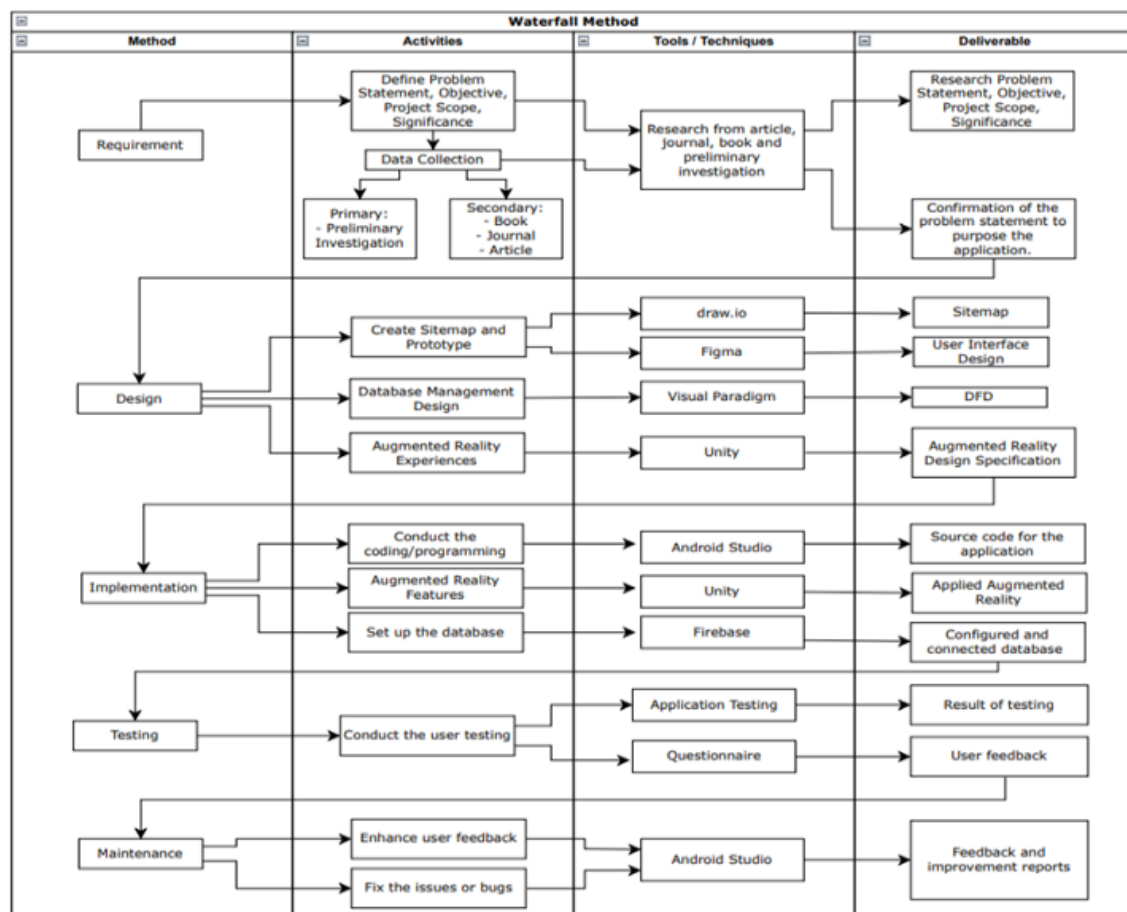


Fig. 1. Methodology, activities, tools, and deliverables of the study

A preliminary investigation has been conducted to understand the initial requirement and need for flower information and AR application use, particularly among students. This information serves as a foundation for further analysis and potential enhancements in flower learning and understanding. The respondents include students in UiTM who are interested in flowers, as well as educators and students involved in planting degree. Out of 107 respondents, 83% showed interest in getting information and knowledge or learning about flowers; 10% had no interest, and 7% had no interest in the process. The respondents were also asked to rate their ability to identify flower species on a scale from 1 to 5, with 1

indicating the least ability and 5 signifying a high ability. The distribution of responses reveals that 21% chose the lowest scale of 1, indicating a limited proficiency in identifying flower species. Most respondents, 37% each, chose scales 2 and 3, suggesting a moderate level of ability. A smaller group of respondents, 3% in total, expressed a relatively high ability by selecting scale 4. Only 2% respondents claimed the highest level of ability, scale 5, in identifying flower species. These findings showed the lack of flowers' information even though the respondents indicated high interest in learning about the flowers. Conventional ways of learning are always limited, which limits their learners in learning (Chamidah et al., 2019).

When asked about the respondents' familiarity with the meaning of flowers and their interest in acquiring information and the knowledge; 32% of respondents indicated that they are familiar with the meaning of flowers. On the other hand, 44% of respondents stated that they are not familiar with the meaning of flowers. Additionally, 24% of respondents expressed uncertainty regarding their familiarity with flower meanings or their interest in obtaining information and knowledge. For question regarding the impact of knowledge on boosting their interest in flowers. The majority of 86% of respondents affirmed that knowledge has a positive influence in increasing their interest in flowers. Conversely, a small minority of 3% respondents indicated that knowledge does not have such an impact, while 11% of respondents expressed uncertainty on this matter.

On the other hand, the investigation also found that 77% of the respondents agreed that technological advancement application, such as AR, can enhance peoples' interest in using for flower learning. 18% were unsure and 5% stated that the technology would not help in the learning. Since AR is still in its infancy in certain country and rural regions that are currently starting to consider the AR use (Avila-Garzon et al., 2021), the question regarding the technology gave clear opportunities of its application in the field. Generally, the findings show the interest in learning about flowers as well as the applications of AR as the information-gathering method. In the design and development phase, Figma, Unity, Android Studio, Visual Code Studio, and Firebase are the main software used to develop the PetalsPalette. Fig. 2 shows the Sitemap for the application. The sitemap created to show and navigate the pages available in the application. The goal of this sitemap is to indicate the flow of the application pages. The sitemap illustrates the application flow for both admin and user, highlighting the differences in navigation and access level between the admin and user applications. The Admin has the access to manage the flower graphics and information that are stored in the database, while the user can utilize the application to view the selected content, including its graphic.

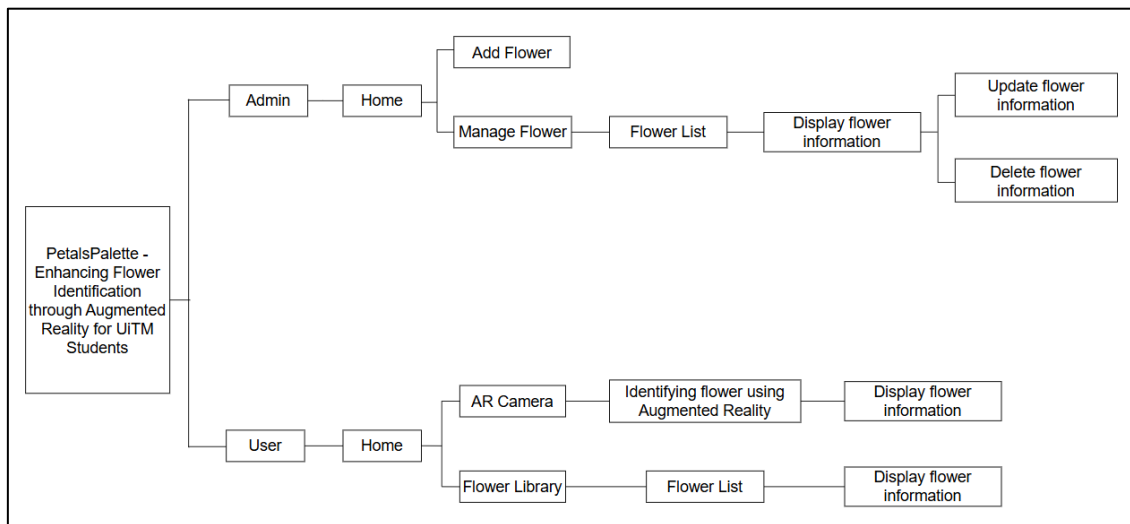


Fig. 2. PetalsPalette sitemap

Fig. 3 shows the flow for DFD of PetalsPalette. Admin sends new flower information to the Adding Flower Data process. Then, the new flower data is stored in the Flower Database. The Flower Database displays flower information to the Admin. After that, the Admin can view the flower information from Flower Database and can manage or update the flower information. Then, Admin sends updated flower information to the Updating Flower Data process and stored it in the Flower Database.

When a user scans a flower marker, the Scanning Marker of Flower process retrieves the flower information from the database in the AR system. Also, the User can view the flower information from the Flower Database. This DFD illustrates the complete cycle of flower data is managed by the Admin and accessible to the User. The Flower Database plays a role in storing, updating, and retrieving the flower data, ensuring both the Admin and the User gets the same flower data.

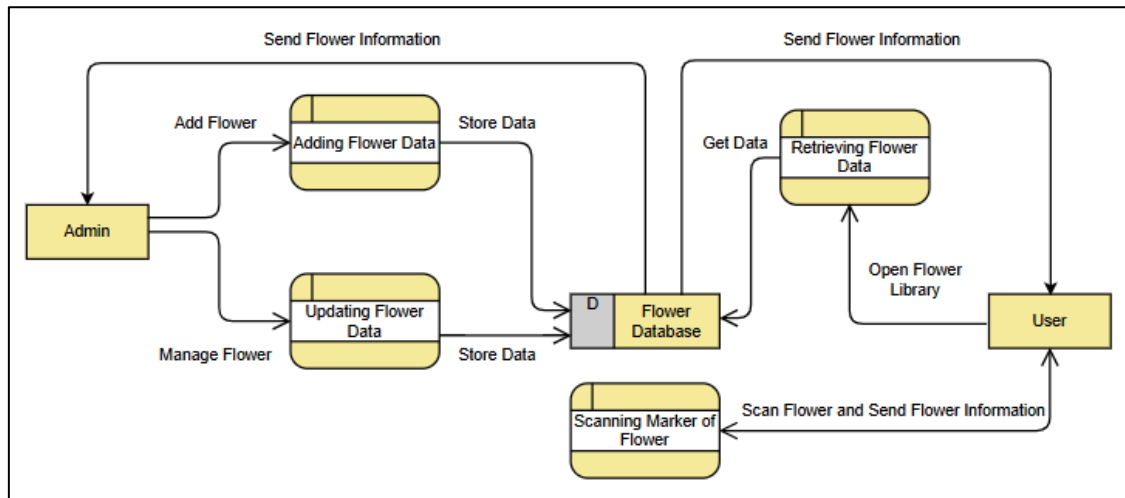


Fig. 3. PetalsPalette data flow diagram

User interface design plays a pivotal role in facilitating visual interaction between users and the application, as demonstrated by the prototypes illustrated. Fig. 4 shows the design sketch and user interface for the admin application, highlighting the functionalities for administrators. The Add Flower page allows the admin to input new flower information into the system. On this page, the admin can upload a flower image, enter the flower's name, and provide a detailed description. The Flower List page provides a comprehensive view of all flowers currently in the system. From this page, the admin can select flowers to update information, modifying the existing details such as the flower image, name, or description. Fig. 5 shows the design sketch and user interface in the user application. There are two buttons for AR Camera and Flower Library. The AR Camera is to enhance the user experience through AR by scanning a marker-based flower image, and it will provide real-time text information about the flower. In Flower Library, the user can view the details about the flower information.

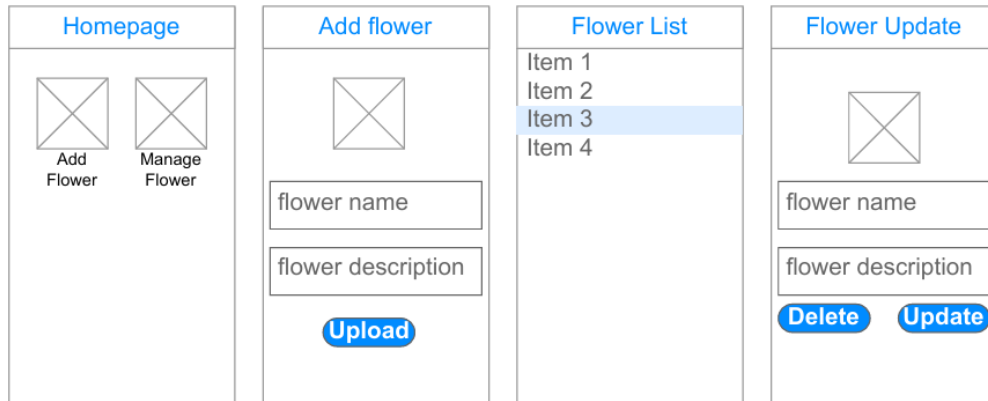


Fig. 4. Admin design sketch

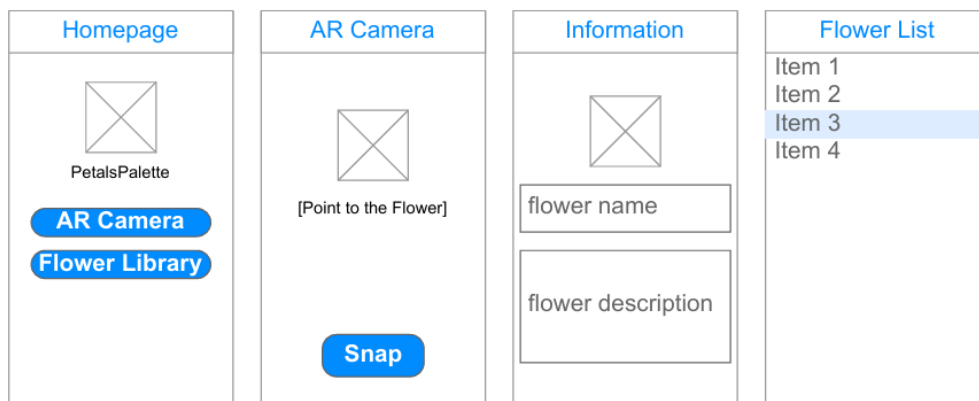


Fig. 5. User design sketch

Firestore was used to store and manage the flower data. Firestore is integrated with Android Studio and Unity, ensuring the data is displayed accurately across both applications. Fig. 6 shows the process of adding flower data within Android Studio through the admin application. The system generates a Flower ID automatically, also known as a unique key. The Flower ID facilitates the insertion of data into Firestore Realtime Database. In Fig. 7, the FlowerDataFetcher.cs is represented within Unity to retrieve data to the Flower Library within the user application. This functionality enables access to flower data stored within application's database, enhancing the user experience by providing timely and accurate information about various flowers that have been added in the database. Realtime Database is used to store the flower data, as shown in Fig. 8.

```

private void uploadToFirebase(Uri uri) {
    String name = flowerName.getText().toString().trim();
    String description = flowerDescription.getText().toString().trim();

    if (!name.isEmpty() && !description.isEmpty()) {
        // Use a more descriptive variable name instead of "imageReference"
        StorageReference flowerImageReference = storageReference.child(pathString: System.currentTimeMillis() + "." + getFileExtension(uri));

        flowerImageReference.putFile(uri).addOnSuccessListener(new OnSuccessListener<UploadTask.TaskSnapshot>() {
            @Override
            public void onSuccess(UploadTask.TaskSnapshot taskSnapshot) {
                flowerImageReference.getDownloadUrl().addOnSuccessListener(new OnSuccessListener<Uri>() {
                    @Override
                    public void onSuccess(Uri uri) {
                        // Generate the key first
                        String key = databaseReference.push().getKey();
                        // Create an instance of NoticeData with flowerName, flowerDescription, and flowerImage
                        NoticeData dataClass = new NoticeData(key, name, description, uri.toString());

                        // Set the value in the database
                        databaseReference.child(key).setValue(dataClass);
                        progressBar.setVisibility(View.INVISIBLE);
                        Toast.makeText(context: AddFlowerPage.this, text: "Uploaded", Toast.LENGTH_SHORT).show();
                        Intent intent = new Intent(packageContext: AddFlowerPage.this, MainActivity.class);
                        startActivity(intent);
                        finish();
                    }
                });
            }
        });
    }
}

```

Fig. 6. Add Flower code

```

D:\> Unity > Projects > PetalPalette > Assets > Script > FlowerDataFetcher.cs
1  using Firebase;
2  using Firebase.Database;
3  using Firebase.Extensions;
4  using UnityEngine;
5  using UnityEngine.UI;
6  using TMPro;
7
8  public class FlowerDataFetcher : MonoBehaviour
9  {
10     public GameObject flowerItemPrefab;
11     public GameObject flowerItemPrefabTransparent;
12     public Transform contentPanel;
13     public ScrollRect scrollRect;
14
15     private DatabaseReference databaseReference;
16
17     void Start()
18     {
19         // Ensure the prefabs are inactive at the start
20         if (flowerItemPrefab.activeInHierarchy)
21         {
22             flowerItemPrefab.SetActive(false);
23         }
24         if (flowerItemPrefabTransparent.activeInHierarchy)
25         {
26             flowerItemPrefabTransparent.SetActive(false);
27         }
28     }

```

Fig. 7. FlowerDataFetcher code

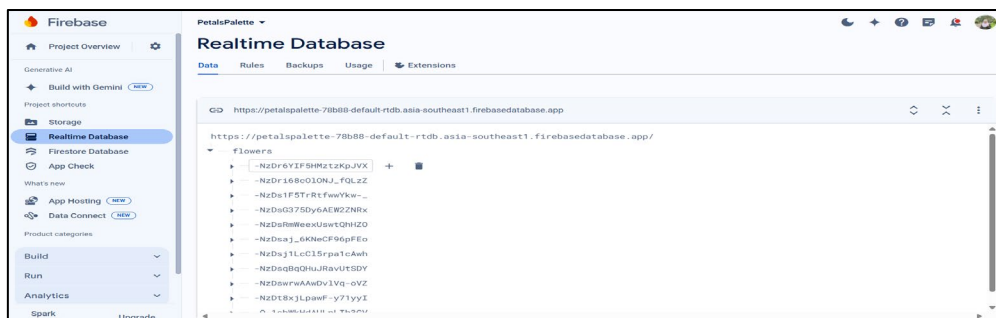


Fig. 8. Realtime database

<https://doi.org/10.24191/jcrinn.v10i2.524>

Vuforia was used to create the image marker for the target image to develop AR features in the PetalsPalette user application. The target image was added into the Vuforia database, it generates the marker features as shown in Fig. 9. The image will display the type of the target, status, target ID, rating of augmentability, and information about when the image was added and modified. After that, all the target images will be downloaded into a database to be exported to Unity. The marker features act to detect the image for AR. Vuforia uses an image marker-based system. The image database is first downloaded as a Unity package. While in Unity, this database is imported into the project. After the import, the image target can be set up by selecting the database from Vuforia for the image target.

Next, as depicted in Fig. 10, the Asset folder within Unity is used to store images of the flower that contain the marker. After setting the image target in Unity, the next step shows in Fig. 11, is to create a canvas that will display the flower name and description when the image marker is detected. This involves designing a user interface (UI) within Unity that includes text fields or panels where the flower name and description will appear. This step ensures that whenever the image marker is recognized by the Vuforia engine, the corresponding information about the flower is presented to the user in a clear and visually appealing manner. This setup enhances the interactive experience by providing relevant details about the flower directly on the screen.

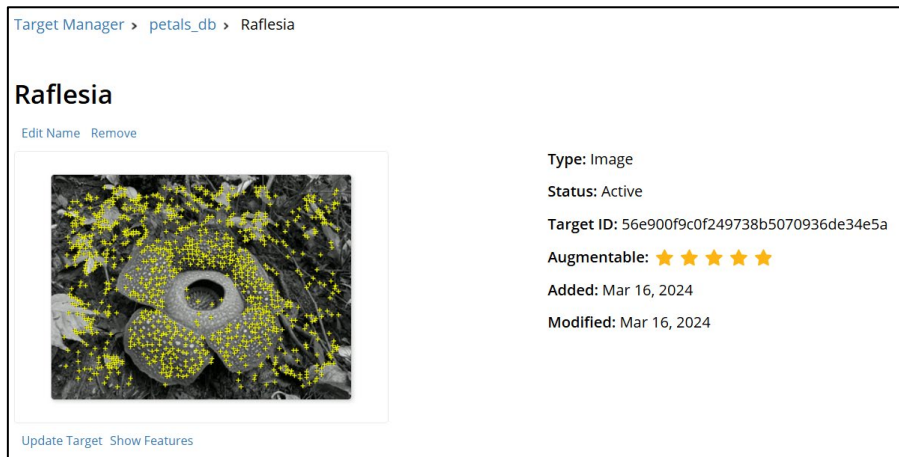


Fig. 9. Flower image marker

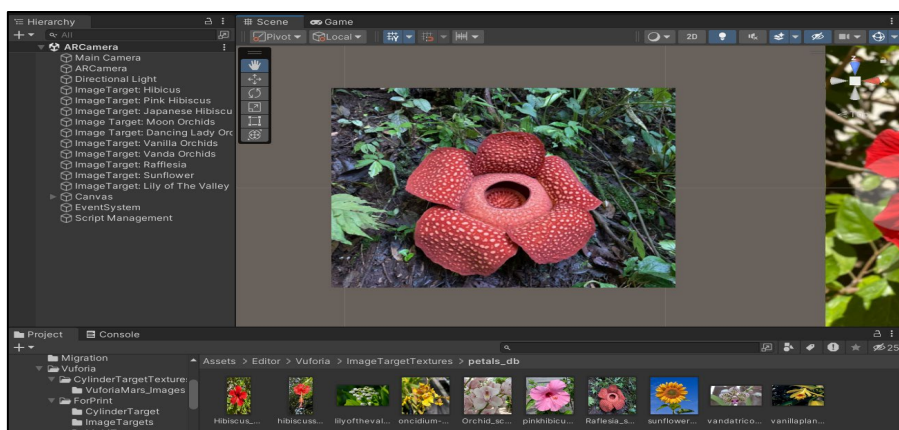


Fig. 10. Image target in Unity

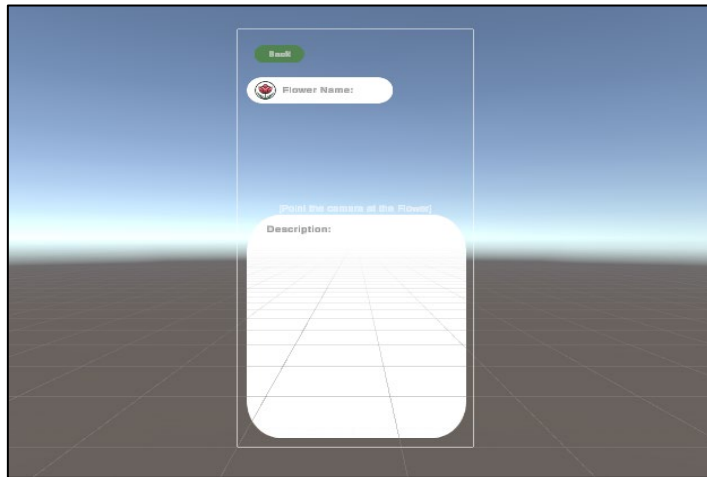


Fig. 11. Image target setting in Unity

Once the flower name and description are set up, they are tracked within the UI, as shown in Fig. 12. The description for each detected flower is attached to the AR trackable element of the image marker. With this setup, when the image marker is detected, the corresponding data that has been inserted for each flower is retrieved and displayed. This ensures that users receive the accurate name and description of the flower in real-time as the AR system recognizes the image markers. Fig. 13 and Fig. 14 show the AR applications of PetalsPalette for admins and users. Fig. 13 enables the admin to add flower information, including its graphics and their characteristics or contents. The added or updated content of the flowers will be saved in the database. Meanwhile, as in Fig. 14, the user is enabled to select and view the content of the flowers. The marker can be scanned, and the content of the selected flowers will be appeared. The user can navigate the content to further understand the flower species.

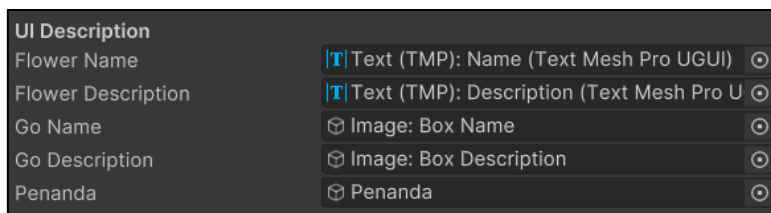


Fig. 12. UI description

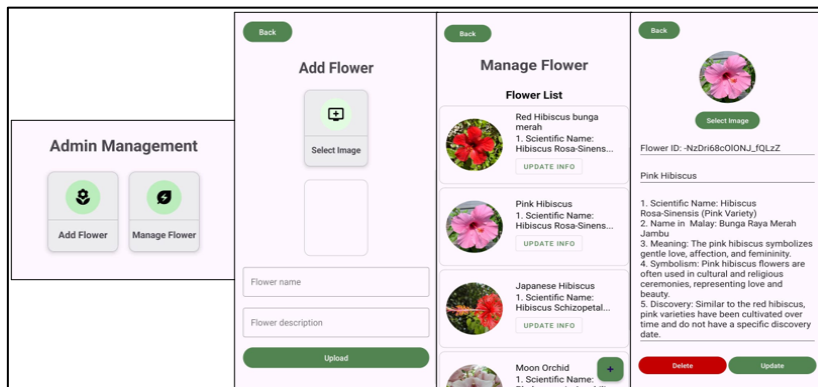


Fig. 13. The developed PetalsPalette (admin)

<https://doi.org/10.24191/jcrinn.v10i2.524>

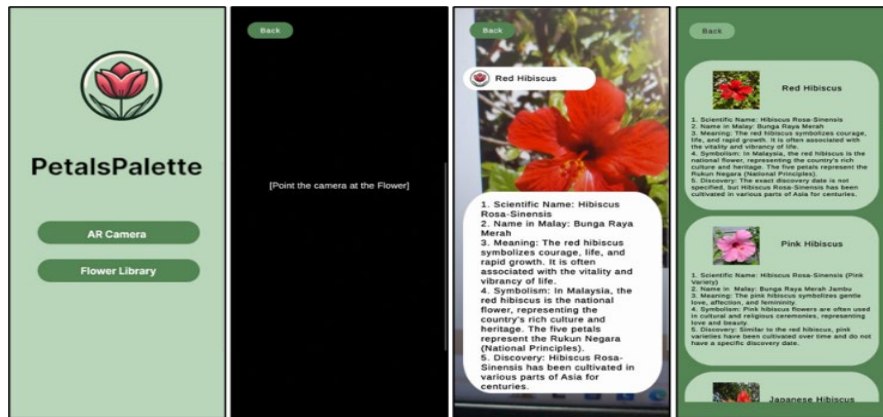


Fig. 14. The developed PetalsPalette (user)

3. PETALSPALETTE TESTING

The Post-Study System Usability Questionnaire (PSSUQ) and User Experience Questionnaire (UEQ) has been used to gather usability and experience perceptions from the respondents. This testing aims to evaluate the usability and user experience of the developed system in identifying the flowers using AR. The testing occurred at an open space in UiTM Arau, Perlis. The goal was to assess if the application could successfully scan the image marker under various lighting conditions.

A guide for users to download the PetalsPalette user application was provided. During the testing, the user received the step-by-step instructions on how to use the application. For the admin application, explanations were given on how the admin could manage the application, with data updates occurring automatically after any changes made by the admin. Users received a Google Form containing PSSUQ and UEQ questionnaires after the testing to provide feedback on the application's usability and functionality. The questionnaires were written in English, included 16 questions from the PSSUQ and 8 questions about the UEQ. A total of 33 UiTM students participated in the user testing, providing valuable responses.

3.1 PSSUQ results

Fig. 15 shows the analysis of usability and experience test results for PSSUQ among UiTM students. In this figure, you can view the mean scores for system usefulness (SYSUSE), information quality (INFOQUAL), and interface quality (INTERQUAL) for the 33 respondents. At the bottom of the figure, the average scores for the mean score, SYSUSE, INFOQUAL, and INTERQUAL are calculated. The average score for the 16 questions from the PSSUQ is under 1.60. This indicates that the respondents generally have a strong agreement with the positive statements about the system's usability, as lower scores on the PSSUQ reflect higher levels of satisfaction and agreement (Mohamad Jamil et al., 2022). Therefore, the users found the PetalsPalette application to be highly useful, with high information and interface quality.

Participant	Questions																Scores on PSSUG			
	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8	Q9	Q10	Q11	Q12	Q13	Q14	Q15	Q16	MEAN SCORE	SYSUSE	INFOQUAL	INTERQUAL
1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1.00	1.00	1.00	1.00
2	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1.00	1.00	1.00	1.00
3	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1.00	1.00	1.00	1.00
4	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1.00	1.00	1.00	1.00
5	1	2	1	1	2	2	1	2	1	1	2	2	2	2	2	2	1.63	1.50	1.50	2.00
6	1	3	2	2	1	3	2	2	1	2	3	1	1	1	1	1	1.69	2.00	1.83	1.00
7	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1.00	1.00	1.00	1.00
8	1	1	1	1	1	1	1	3	2	1	1	1	1	1	1	1	1.19	1.00	1.50	1.00
9	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1.00	1.00	1.00	1.00
10	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1.00	1.00	1.00	1.00
11	1	1	1	1	1	1	2	2	2	1	1	2	2	1	2	2	1.44	1.17	1.67	1.67
12	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1.00	1.00	1.00	1.00
13	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1.00	1.00	1.00	1.00
14	1	1	2	2	2	2	2	2	2	2	1	2	1	2	1	1	1.69	1.67	1.83	1.67
15	1	1	1	1	1	1	2	7	7	2	2	2	2	2	2	1	2.25	1.17	3.67	2.00
16	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1.00	1.00	1.00	1.00
17	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1.00	1.00	1.00	1.00
18	1	1	1	1	1	1	1	5	1	2	1	1	2	2	2	1	1.50	1.00	2.00	1.67
19	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1.00	1.00	1.00	1.00
20	1	1	1	1	1	1	2	2	2	2	1	2	2	1	1	2	1.44	1.17	1.83	1.00
21	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1.00	1.00	1.00	1.00
22	1	2	2	2	2	2	2	2	2	2	2	2	2	2	2	2	1.94	1.83	2.00	2.00
23	1	1	2	2	2	2	4	2	3	1	1	2	1	2	3	1	1.88	1.67	2.17	2.00
24	1	1	1	1	1	1	2	2	3	2	2	1	1	2	2	1	1.50	1.17	1.83	1.67
25	1	1	1	2	2	1	2	2	3	3	2	3	3	3	4	2	2.19	1.33	2.50	3.33
26	1	2	2	2	2	2	2	1	2	1	1	1	3	2	1	1	1.63	1.83	1.33	2.00
27	1	1	1	1	1	1	1	1	2	1	1	1	1	1	1	1	1.06	1.00	1.17	1.00
28	1	1	1	1	1	1	1	5	5	3	3	1	1	1	1	2	1.81	1.00	3.00	1.33
29	1	2	1	2	2	2	1	2	1	1	2	2	2	2	2	2	1.69	1.67	1.50	2.00
30	1	2	2	2	2	2	2	2	2	2	2	2	2	2	2	2	1.94	1.83	2.00	2.00
31	1	2	1	2	2	3	1	2	2	1	2	1	2	2	2	1	1.69	1.83	1.50	2.00
32	1	2	2	2	2	2	2	2	2	1	2	1	1	1	2	1	1.63	1.83	1.67	1.33
33	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1.00	1.00	1.00	1.00
Average Score																	1.39	1.26	1.53	1.41

Fig. 15. PSSUG calculation and results

3.2 UEQ results

Fig. 16 shows the results of the UEQ. It displays the mean scores based on the questions answered by the participants. The mean score for UEQ questions is about 4.5. This value indicates that the respondents generally have a strong positive agreement with the statements about their user experience while using the PetalsPalette application. Since the UEQ score ranges from 1 to 5, with 5 indicating strong agreement, the users found the application provides a fantastic user experience.

Participant	Questions							
	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8
1	5	5	5	5	4	3	4	4
2	5	5	5	5	5	4	5	5
3	5	5	5	5	5	5	5	5
4	5	5	5	5	5	5	5	5
5	5	5	4	5	5	5	5	4
6	5	5	5	5	5	5	5	5
7	5	5	5	5	5	5	5	5
8	5	5	5	5	5	5	5	5
9	5	5	5	5	5	5	5	5
10	5	5	5	5	5	5	5	5
11	5	5	5	5	4	5	4	5
12	5	5	5	5	5	5	5	5
13	5	5	5	5	5	5	5	5
14	5	5	5	5	5	5	5	5
15	5	5	5	5	5	5	5	5
16	5	5	5	5	5	5	5	5
17	5	5	5	5	5	5	5	5
18	5	5	5	5	5	4	5	5
19	5	5	5	5	5	5	5	5
20	5	5	5	5	5	5	5	5
21	5	4	5	5	5	5	5	5
22	5	4	4	4	4	4	4	4
23	5	5	4	5	4	4	5	5
24	5	5	5	5	4	4	5	4
25	5	5	4	5	4	4	5	5
26	5	5	5	5	5	4	5	4
27	5	5	5	5	5	4	5	5
28	5	5	5	5	3	4	4	4
29	5	4	4	5	4	4	4	4
30	5	4	4	4	4	4	4	4
31	5	3	4	5	4	4	4	4
32	5	4	4	5	4	5	5	5
33	5	5	4	5	5	5	4	5
Mean Score	5	4.79	4.73	4.94	4.64	4.58	4.76	4.73

Fig. 16. UEQ calculation and results

4. DISCUSSION

The overall findings for PSSUQ and UEQ collected the positive responses from the respondents, indicating that the PetalsPalette application effectively enhances flower identification through AR. Table 1 shows the overall PSSUQ results: Overall Satisfaction has a mean score of 1.39, SYSUSE (System Usefulness) has a mean score of 1.26, INFOQUAL (Information Quality) has a mean score of 1.53, and INTERQUAL (Interface Quality) has a mean score of 1.41. These low scores indicate that the respondents were highly satisfied with the application's usability, usefulness, information quality, and interface quality. Previous studies in AR for learning also showed that the PSSUQ score indicates that the AR application is well-received, with the score for each dimension are approaching 1 (1 indicates a strong agreement), but opportunities for enhancing the user experience persist (Kheng, 2023).

Table 2 shows the overall UEQ results for the 8 questions. The first question, which assesses attractiveness, received a score of 5. The second and third questions, which evaluate perspicuity, both received a score of 4.76. The fourth question, related to efficiency, scored 4.94. The fifth question, which measures dependability, scored 4.64. The sixth and seventh questions, which pertain to stimulation, both scored 4.67. The final question, the eighth, which assesses novelty, received a score of 4.73. The average user experience (UX) result is 3.59. These results indicate that respondents found the PetalsPalette application to be highly attractive, easy to understand (perspicuity), efficient, dependable, stimulating, and novel. The overall positive scores across different user experience dimensions demonstrate that the application provides a favorable user experience. in terms of the feelings and impressions made while interacting with the developed application. This aligns with the study by Derisma and Hersyah (2021), which reported positive findings regarding the evaluation of AR usage in learning. The utilization of AR provides significant feelings of learning, impressions, and positive attitudes towards the learning content. In their study, Ramli et al. (2023) also show positive findings of UEQ in the utilization AR and game elements in STEM content. Studying STEM subjects in an online medium, as well as conventional methods, might be hard and challenging, with technical vocabulary, intricate images, and complex explanations. With AR, the learning experience can be more engaging and benefit the learners.

As for AR marker detection from various lighting, the testing showed that lighting did affect the marker detection. The lighting made the content disappear or fail to appear in the session. However, finding suitable lighting in the testing location solved the problem easily.

Table 1. PSSUQ results

Scores on PSSUQ	Mean Score; Range 1–7
Overall Satisfaction	1.39
System Usefulness (SYSUSE)	1.26
Information Quality (INFOQUAL)	1.53
Interface Quality (INTERQUAL)	1.41

Table 2. PSSUQ results

Questions	Aspect	Average Answer Result
Q1	Attractiveness	5
Q2	Perspicuity	4.76
Q3		
Q4		
Q5	Dependability	4.64
Q6	Stimulation	4.67
Q7		
Q8	Novelty	4.73
UX Average Result:		3.59

5. CONCLUSION

This paper offers an in-depth primer on a study called PetalsPalette, which highlights flower learning through AR for students. The study concludes that the users have successfully embraced the PetalsPalette mobile application due to its robust usability and experience. Despite these achievements, limitations were acknowledged, including the exclusive focus on UiTM students and constraints in AR technology, such as static text updates. Suggestions for future work include reaching out to a wider range of users beyond UiTM, testing the app on different devices and operating systems, and improving AR technology to make the application more flexible and user-friendly for everyone. Broader participation ensures the generalization of the findings and can show the acceptability of the proposed application. These improvements aim to advance AR technology in flower identification and enhance overall usability across diverse user groups and technological environments.

6. ACKNOWLEDGEMENTS/FUNDING

The authors would like to acknowledge the support of the Faculty of Computer Science and Mathematical Studies, Universiti Teknologi Mara (UiTM), Cawangan Perlis, Kampus Arau for approving this study. Much appreciation also goes to all the involved personnel and individuals throughout the study.

7. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts.

8. AUTHORS' CONTRIBUTIONS

Nur Azmatun Farwizah carried out, analysed, and wrote the research report; **Nurtihah Mohamed Noor** supervised, guided, and provided solutions throughout the research duration; **Mohd. Fitri Yusoff** offered ideas and solutions, organized, and edited this article.

9. REFERENCES

- Avila-Garzon, C., Bacca-Acosta, J., Duarte, J., & Betancourt, J. (2021). Augmented reality in education: An overview of twenty-five years of research. *Contemporary Educational Technology*, 13(3), ep302. <https://doi.org/10.30935/cedtech/10865>
- Chamidah, D., Kristianto, S., Sunaryo, Fajarianto, O., Ahmad, A., Ani Setyo Dewi, Y., Sambodja, E., Ambarumi Munawaroh, D., Fitriah, N., & Indriawati, P. (2019). Feasibility of based augmented reality devices discovery learning on students learning outcomes in morphology of Wijaya Kusuma flower epiphyllum anguliger. *Journal of Physics: Conference Series*, 1175, 012261. <https://doi.org/10.1088/1742-6596/1175/1/012261>
- Derisma, D., & Hersyah, M. (2021). User experience measurement using augmented reality application in learning 4.0. In *ICED-QA 2019: Proceedings of the 2nd International Conference on Educational Development and Quality Assurance* (pp. 230-240). European Alliance for Innovation.
- Kheng, T. K. (2023). *Interactive augmented reality aided software for Physics education using exploratory approach*. <http://eprints.utar.edu.my/6102/>
- Kim, I., Park, J.-S., & Choi, K.-O. (2020). Analysis of designation and symbolic meanings of floral emblems in South Korea as elements of garden tourism and design. *Journal of People, Plants, and Environment*, 23(1), 87–99. <https://doi.org/10.11628/ksppe.2020.23.1.87>
- Lam, K. Y., Lee, L. H., & Hui, P. (2021). A2W: Context-aware recommendation system for mobile augmented reality web browser. In *Proceedings of the 29th ACM International Conference on Multimedia* (pp.2447–2455). ACM. <https://doi.org/10.1145/3474085.3475413>
- Łysiak, G. (2022). Ornamental flowers grown in human surroundings as a source of anthocyanins with high anti-inflammatory properties. *Foods*, 11(7), 948. <https://doi.org/10.3390/foods11070948>
- Marchenko, A. M., & Kuzovkina, Y. A. (2021). Calculation of the ovule number in the genus salix: A method for taxa differentiation. *Applications in Plant Sciences*, 9(11–12). <https://doi.org/10.1002/aps3.11450>
- Moesl, B., Schaffernak, H., Vorraber, W., Holy, M., Herrele, T., Braunstingl, R., & Koglbauer, I. V. (2022). Towards a more socially sustainable advanced pilot training by integrating wearable augmented reality devices. *Sustainability*, 14(4), 2220. <https://doi.org/10.3390/su14042220>
- Mohamad Jamil, P. A. S., Karuppiah, K., Mohammad Yusof, N. A. D., Mohd Suadi Nata, D. H., Abdul Aziz, N., How, V., Mohd Tamrin, S. B., & Naeni, H. S. (2022). Usability testing of a wireless individual indicator system application: monitoring exposure to outdoor air pollution among Malaysian traffic police. *Digital Health*, 8, 205520762211033. <https://doi.org/10.1177/20552076221103336>
- Narbona, E., Arista, M., Whittall, J. B., Camargo, M. G. G., & Shrestha, M. (2021). Editorial: The role of flower color in angiosperm evolution. *Frontiers in Plant Science*, 12. <https://doi.org/10.3389/fpls.2021.736998>
- Parsons, T. D., Gaggioli, A., & Riva, G. (2020). Extended reality for the clinical, affective, and social

- neurosciences. *Brain Sciences*, 10(12), 922. <https://doi.org/10.3390/brainsci10120922>
- Rajesh Kumar, K., & Subhashini, S. J. (2023). Visualizing medical flowers details by using deep neural network. *Recent Developments in Electronics and Communication Systems*, 208-215. <https://doi.org/10.3233/ATDE221259>
- Ramli, R. Z., Sahari Ashaari, N., Mat Noor, S. F., Noor, M. M., Yadegaridehkordi, E., Abd Majid, N. A., ... & Abdul Wahab, A. N. (2024). Designing a mobile learning application model by integrating augmented reality and game elements to improve student learning experience. *Education and Information Technologies*, 29(2), 1981-2008. <https://doi.org/10.1007/s10639-023-11874-7>
- Setyaningsih, Y., Dijaya, R., & Suprianto, S. (2021). Ethnoscience based augmented reality on botanical garden. *JUITA: Jurnal Informatika*, 9(2), 173. <https://doi.org/10.30595/juita.v9i2.10602>
- Tonmoy, T. K., & Islam, Md. A. (2023). Do students look for information differently? Information-seeking behavior during the Covid-19 pandemic. *Digital Library Perspectives*, 39(2), 166–180. <https://doi.org/10.1108/DLP-09-2022-0073>
- Vidak, A., Movre Šapić, I., & Mešić, V. (2022). Augmented reality in teaching about Physics: First findings from a systematic review. *Journal of Physics: Conference Series*, 2415(1), 012008. <https://doi.org/10.1088/1742-6596/2415/1/012008>
- Wilujeng, S., Chamidah, D., & Wahyuningtyas, E. (2019). The use of augmented reality to introduce Wijaya Kusuma flower. In *Proceedings of the 6th International Conference on Community Development*. <https://doi.org/10.2991/iccd-19.2019.135>
- Wong, J., Bayoumy, S., Freke, A., & Cabo, A. (2022). Augmented reality for learning Mathematics: A pilot study with WebXR as an accessible tool. In *Towards a New Future in Engineering Education, New Scenarios That European Alliances of Tech Universities Open Up* (pp. 1805–1814). <https://doi.org/10.5821/conference-9788412322262.1216>
- Wulansari, T. Y. I., Senjaya, S. K., & Astuti, I. P. (2022). Flower structures of averrhoa dolichocarpa Rugayah & Sunarti. *Journal of Tropical Biodiversity and Biotechnology*, 7(3), 74585. <https://doi.org/10.22146/jtbb.74585>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Assessing Knowledge Acquisition and Perceptions in Hands-On Video Editing Workshop for Non-Technical Students

Zubaidah Bohari¹, Abdul Hadi Abdul Talip^{2*}, Satria Arjuna Julaihi³, Rumaizah Che Md Nor⁴, Norizuandi Ibrahim⁵

^{1,2,3,4,5}Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Sarawak Branch, Samarahan Campus, 94300 Kota Samarahan, Sarawak, Malaysia.

ARTICLE INFO

Article history:

Received 15 May 2025

Revised 15 June 2025

Accepted 25 June 2025

Online first

Published 1 September 2025

Keywords:

Video Editing

Knowledge Acquisition

Montage Elements

Hands-On Learning

Non-Technical Students

DOI:

10.24191/jcrinn.v10i2.513

ABSTRACT

In the era of digital technology, video editing is now a priority skill in various fields like education, marketing, and business. Nevertheless, non-technical students tend to lack video editing skills because of their limited exposure to montage elements and video editing software. This research explores the effects of practical video editing workshops on knowledge and perceptions of non-technical students. The workshops were aimed at imparting knowledge of montage elements, a key component of video production. With a quantitative research design, 30 students from non-technical disciplines have participated in guided workshops with Microsoft PowerPoint and CapCut. The findings indicate a significant improvement in the comprehension of montage elements, including visual organization, audio synchronization, transitions, and pacing among the students, with high mean knowledge scores (3.97 - 4.05) on a 5-point scale. In addition, regression analysis shows a significant relationship ($p\text{-value} < 0.05$), confirming its effectiveness in engaging non-technical students and achieving a better understanding of video editing concepts. The results of this research contribute to the development of instructional strategies that promote improved learning experiences for groups of students with varying backgrounds.

1. INTRODUCTION

In an increasingly digital world, the necessity for video editing skills has extensively spread across many disciplines, including education, marketing, and business, in addition to its initiation in media studies (Brown & Green, 2021). Video production is no longer confined to professional filmmakers but has become an essential skill for educators, students, and content creators across industries (Snelson, 2018). However, mastering the key montage elements, such as visual, audio, text, graphic, transition, and pacing, remains a challenge for non-technical students due to limited exposure to video editing techniques and editing software. Moreover, many traditional classroom approaches focus on theoretical knowledge rather than

^{2*} Corresponding author. E-mail address: adie0951@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.513>

hands-on application, limiting students' ability to develop practical skills (AlKhaibary et al., 2020). Without adequate exposure to these elements, non-technical students may struggle with video composition and effective content delivery. It is important to understand these elements as they help drive the overall narration, coherence, and appeal of a video project (Wan et al., 2022). Integrating digital literacy into higher education curricula supports student readiness for the digital workplace and helps bridge competency gaps (Kayyali, 2024). Additionally, positive attitudes toward digital technology significantly predict students' confidence and engagement in online learning environments (Getenet et al., 2023).

Hands-on workshops provide an experiential learning opportunity where students learn technical and creative skills by doing things under supervision. According to experiential learning theories, active participation in supervised tasks enhances conceptual knowledge and skill development (Kolb, 1984). Therefore, this study focuses on addressing two key objectives: (1) To evaluate the knowledge level of montage elements after attending the Hands-On Video Editing Workshops, and (2) To evaluate the relationship between students' knowledge of montage elements and their perceptions of the Hands-On Video Editing Workshops.

This research also aims to provide insights into teaching instructional strategies that can improve learning outcomes among non-technical students by evaluating the effectiveness of practical video editing workshops. Understanding how students view these workshops and how they affect their understanding of montage elements will help create more engaging and effective multimedia teaching and learning experiences.

2. LITERATURE REVIEW

Evaluating the impact of hands-on video editing workshops on participants' understanding of montage elements is crucial for assessing the effectiveness of such educational programs. The field of video editing has long been shaped by both theoretical frameworks and practical innovations. Dancyger (2018) highlights the transition from linear analogue editing systems to non-linear digital platforms, underscoring the transformative impact of software like Adobe Premiere Pro and Final Cut Pro on creative workflows. He argues that these advancements have not only increased efficiency but also expanded the creative possibilities for editors, enabling them to experiment with complex narratives and visual styles. This literature review explores the significance of montage elements in video editing, the benefits of hands-on workshops, and methods for evaluating knowledge acquisition in this context. Besides, it also explores how students' knowledge and perceptions interact in the context of hands-on video editing workshops, drawing from established educational theories and empirical studies.

2.1 The role of montage elements in video editing

A montage is defined as "the production of a rapid succession of images in a movie to illustrate an association of ideas" (Merriam-Webster, 2025). According to Xiang et al. (2022), montage is one of the vital components of a film, where it can create a story by rearranging images and shots, and it can even illustrate the essence of ideas to the viewer. They further explained that this technique is used to shorten time, shows growth in a character and even expresses complex information effectively. Understanding the elements is crucial for beginners learning video editing, as these elements form the building blocks of effective storytelling. Each montage is unique, and not all montages incorporate every element. However, certain elements are commonly used to enhance the storytelling impact.

These elements include music, quick cuts, voiceover narration, minimal or no dialogue, and repeated camera movements (DeGuzman, 2024). Music plays a crucial role in setting the tone and pacing, guiding the audience's emotional response. Quick cuts maintain energy and urgency by rapidly transitioning between shots, while voiceover narration provides additional context or insight to reinforce the montage's

message. In many cases, minimal or no dialogue is used, emphasizing visual storytelling through expressive imagery and action. Additionally, repeated camera movements create a sense of rhythm and cohesion, making the montage feel seamless and engaging. Together, these elements enhance storytelling impact, ensuring the montage effectively conveys its intended message. For non-technical students learning video editing, understanding these elements is particularly important. Hands-on workshops that focus on applying montage techniques can help students grasp the interplay between visuals, audio, and pacing, enabling them to create cohesive and engaging videos.

In educational contexts, such workshops have proven effective in teaching montage techniques to beginners, enabling them to develop both technical skills and creative confidence. Montage elements play a vital role in video editing by enhancing storytelling and creating engaging visuals. Understanding these elements is essential for beginners, particularly non-technical students, who aim to develop video editing skills. The hands-on video editing workshop in this study focuses on teaching these elements, providing students with the tools they need to create effective montages. The following sections will explore how knowledge of montage elements influences students' perceptions of the workshop.

2.2 Impact of hands-on lab on non-technical students

Hands-on labs are a transformational way to teach technical skills to students without a background in technology. For instance, in this research, the students involved are from programs like Public Administration or Halal Management at UiTM Sarawak and have little to no experience with video editing tools. A study done by Kirk (2015) highlighted how student-centered hands-on labs bridge this gap by allowing students to learn through practice. He noted that the "learn by doing" approach helps non-technical students gain confidence faster than traditional lectures. In the context of this study, students used tools like CapCut and PowerPoint to edit videos, which aligns with the modern approach to experiential learning.

One major advantage of hands-on labs is that they make abstract ideas clearer. For example, when learning montage elements like pacing or transitions, students might struggle to understand these concepts in theory. However, when they received hands-on instruction to create a short video, the concepts became easier to grasp. Research by Zhang et al. (2022) highlights that students taking hands-on activities often thrive when they see their work come to life, such as editing a montage for a class project. In certain cases, students can feel overwhelmed by tools like CapCut at first. To help them, tasks were broken down into smaller steps, with step-by-step guidance, a strategy Theobald et al. (2020) call "scaffolding," which reduces frustration and builds confidence.

As shown by Knoblauch (2022), students who went through experiential learning with a focus on the experience had gained new skills and knowledge. He reported that this approach had helped the students immensely. In the case of this study, the hands-on lab focused on practical tasks, like arranging visuals or syncing music, thus making sure students can develop their creativity while learning technical skills.

2.3 Benefits of hands-on video editing workshops

Immersive learning experiences that connect theoretical knowledge with practical application are offered by practical, hands-on workshops. Engaging directly with editing software and projects enables participants to develop both technical proficiency and creative decision-making skills. Such workshops often cover essential aspects of video editing, including technical proficiency, creative techniques and critical analysis.

The first aspect, technical proficiency, involves participants gaining a solid foundation for professional editing by learning how to use editing software, control timelines, and apply effects. The second aspect, creative techniques, oversees how a workshop emphasizes the artistic aspects of editing, such as storytelling, pacing, and the effective use of montage to convey complex ideas succinctly. Finally, critical

analysis encourages learners to critique their own work and that of others, fostering a deeper understanding of editing choices and their impact on audiences.

Dorfman et al. (2019) highlight the importance of personalized multimedia content in educational settings. Their study, titled “Teachers personalize videos and animations of biochemical processes: results from a professional development workshop”, demonstrates how hands-on workshops empower educators to customize educational videos and animations to better suit their students' learning needs. This personalization not only enhances the clarity and effectiveness of instructional materials but also helps teachers develop their technical proficiency and creative confidence in multimedia production.

Collaboration and industry networking are also key benefits of video editing workshops. Jeong and Lim (2024), in their paper “Implementing a video production collaboration system with information marker overlay technique”, emphasize the role of collaborative systems in enhancing team-based video production. Their study illustrates how collaborative environments foster efficient communication, idea sharing, and division of labor, ultimately leading to more polished and innovative multimedia content. The integration of industry-standard tools and techniques further provides participants with networking opportunities, connecting them with professionals and potential collaborators in the field.

Career development and portfolio building are additional crucial outcomes of hands-on video editing workshops. Yen and Yang (2024), in their study “The Integration of Artificial Intelligence and Video Production Skills in Workplace Development: A Study from the Perspective of Vocational Training”, discuss how practical video editing skills, combined with emerging technologies like artificial intelligence, are becoming essential for workplace readiness. Workshops that emphasize real-world projects enable participants to create high-quality portfolios that showcase their technical capabilities and creative vision. These portfolios serve as vital tools for job applications and career advancement, helping participants stand out in competitive job markets.

The role of video in education is another important dimension of video editing workshops. Priandika et al. (2022), in their paper “Video Editing Training to Improve the Quality of Teaching and Learning at SMK Palapa Bandarlampung”, demonstrate how video editing training enhances the quality of educational content. Their study shows that teachers who participate in video editing workshops produce more engaging and effective instructional materials, leading to improved student understanding and participation. By equipping educators with the skills to create customized, high-quality videos, these workshops contribute directly to enhancing the teaching and learning experience.

2.4 Utilizing online questionnaires for assessing knowledge and skills in hands-on workshops

Online questionnaires have become a popular tool for evaluating student learning in hands-on workshops, especially in technical fields like video editing. These surveys are practical because they allow immediate feedback collection after an activity, ensuring good and effective responses that can help future planning (So & Chan, 2018). For non-technical students, online questionnaires are user-friendly, as they can be completed on smartphones or laptops, aligning with how Malaysian students typically interact with technology in their daily lives.

A key advantage of online questionnaires is their ability to measure knowledge (e.g., understanding of montage elements) and skills (e.g., ability to use editing tools like MS PowerPoint and CapCut). A study by Sufi et al. (2018) mentioned that the usage of a questionnaire is to measure the impact of a workshop. They found that questionnaires effectively captured improvements in both technical knowledge and practical abilities. Similarly, Scheef and Johnson (2017) highlighted that online tools like Google Forms or SurveyMonkey simplify data analysis, allowing educators to quickly identify areas where students struggle, such as syncing audio or applying transitions.

2.5 Student perceptions of hands-on video editing workshops

Anas (2019) suggests that student-created video projects enhance active learning by encouraging students to take a more engaged role in their education. This reinforces the idea that video projects are both enjoyable and beneficial for students. Perceptions about hands-on learning experiences are also crucial to understanding the effectiveness of video editing workshops. Nelson (2018), in their doctoral dissertation “Perceptions about Hands-On Art Making by Online Students”, explores how students engaged in hands-on art-making activities perceive their learning experiences compared to virtual-only instruction. The findings indicate that students value the tactile and interactive elements of hands-on work, which contribute to a greater sense of engagement, creativity, and satisfaction. This aligns with the principles of video editing workshops, where hands-on practice enhances both technical proficiency and creative expression.

2.6 Knowledge acquisition through hands-on video editing workshops

Hands-on learning approaches, particularly in video editing, have been shown to enhance students' technical skills and conceptual understanding. Engaging directly with video production tasks allows students to apply theoretical knowledge in practical settings, fostering deeper learning. For instance, a study by Galatsopoulou et al. (2022) found that students who participated in video-based active learning scenarios demonstrated improved comprehension of course material. The interactive nature of these workshops encourages learners to experiment with editing techniques, leading to a more robust grasp of the subject matter.

Knowledge acquisition through hands-on video editing workshops is further supported by Fleischmann (2021) in their study “Hands-on versus virtual: Reshaping the design classroom with blended learning”. Fleischmann argues that hands-on learning environments provide more effective knowledge retention and skill development compared to purely virtual formats. By engaging in practical, real-world tasks, participants gain a comprehensive understanding of video editing tools and techniques, while also developing critical problem-solving and creative thinking abilities. This hands-on approach ensures learners are better prepared to apply their knowledge in professional and educational settings.

2.7 Relationship between knowledge acquisition and student perceptions

The relationship between knowledge acquisition and student perceptions in video editing workshops is also a critical area of exploration. Aggad (2023), in their doctoral dissertation “Challenges, Opportunities, and Success Factors of Integrating Video Production Education in Lebanon”, investigates how students' perceptions of video production education influence their knowledge acquisition and engagement. The study reveals that students who perceive video editing workshops as interactive and practically valuable tend to demonstrate higher levels of motivation and a deeper understanding of multimedia concepts. This connection between positive perceptions and enhanced learning outcomes underscores the importance of designing workshops that balance technical training with creative exploration and collaborative opportunities.

The interplay between knowledge acquisition and student perceptions is evident in several studies. Positive perceptions of hands-on activities often correlate with enhanced learning outcomes. For example, Galatsopoulou et al. (2022) found that students' positive attitudes towards video-based learning were significantly related to their intention to engage with the material, which in turn improved their knowledge acquisition. This suggests that when students find learning activities enjoyable and relevant, they are more likely to invest effort, leading to better understanding and retention of information. For instance, Hsieh and Knudson (2020) found that active learning strategies in undergraduate biomechanics courses positively influenced students' learning outcomes. It shows that there is a significant correlation between students' positive perceptions of hands-on learning activities and enhanced learning outcomes.

3. METHODOLOGY

This study used a quantitative research design with a survey approach to assess students' understanding of montage elements and a correlational design to investigate the relationship between their knowledge and perceptions toward the hands-on video editing workshop. The study involved 30 students enrolled in the "Computer and Information Processing" course during the March-October 2024 semester, from the Diploma in Public Administration (AM110) and the Diploma in Halal Management (IC120) at Universiti Teknologi MARA Sarawak, Kampus Samarahan 2. These students possessed no formal education in video editing or computer science. Questionnaires were completed by 27 out of the 30 students and a response rate of 90% was attained.

The "Video Montage Workshop: Hands-on Editing" was held in two sessions. The main aim was to develop video editing skills using widely available software tools, Microsoft PowerPoint and CapCut. The workshop aimed to give a hands-on experience. It had two main parts: "Guided Practical Sessions," which taught how to use the elements of montage in video editing through step-by-step instructions, and "Introduction to Video Editing Software," introducing the participants to the interface and the basics.

Data collection was conducted using online questionnaires administered immediately after the workshop ended. The questionnaires were designed to evaluate students' video editing skills and assess their knowledge of video editing software.

The survey data was analysed using descriptive statistics, correlation analysis, and regression analysis. Descriptive statistics summarised the data by calculating measures such as means and frequencies. Correlation analysis investigated the association between students' knowledge and their perceptions of the workshop, whereas regression analysis assessed the significance of this association by assessing how much knowledge influenced students' perceptions. Fig. 1 shows the overall workflow of the research methodology, including the workshop implementation, data collection approach, and subsequent data analysis.

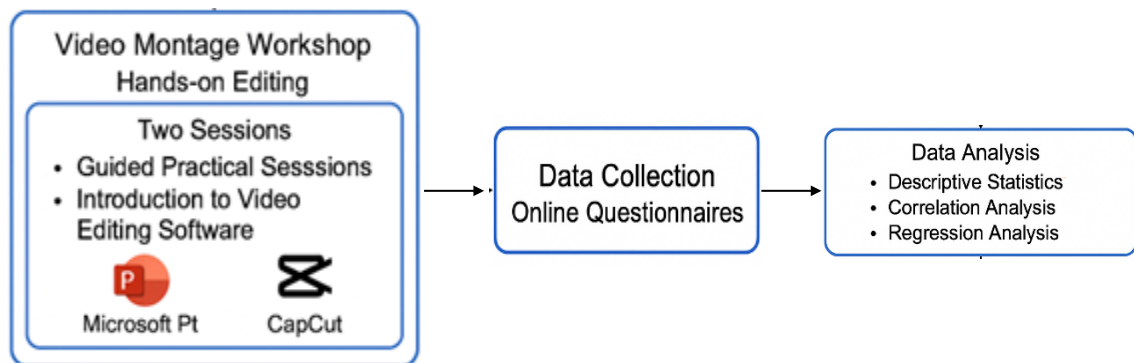


Fig. 1. Research methodology workflow

4. ANALYSIS AND DISCUSSION

4.1 Demographic profile

Table 1. Demographics of respondents

Demographic	Label	Frequency	Percentage (%)
Gender	Male	3	11.1
	Female	24	88.9
Semester of study	1	13	48.1
	2	14	51.9
Faculty	Academy of Contemporary Islamic Studies	14	51.9
	Faculty of Administrative Science and Policy Studies	13	48.1
Course	Diploma in Halal Management	14	51.9
	Diploma in Public Administration	13	48.1

Table 1 summarizes the demographic profile of the respondents according to gender, semester of study, faculty, and program. Most of the respondents were female students, with 24 students (88.9%), while only 3 male students (11.1%) were involved in this study. Out of 27 respondents, 13 respondents were semester 1 students (48.1%) and the remaining 14 respondents were from semester 2 (51.9%). Students in Semester 1 pursuing a Diploma in Public Administration from the Faculty of Administrative Science and Policy Studies. Additionally, 14 students from semester 2 were enrolled in the Diploma in Halal Management program at the Academy of Contemporary Islamic Studies.

4.2 Evaluate the knowledge level of montage elements after attending the hands-on video editing workshops

All respondents were asked to rate the knowledge level of the montage elements used during the workshop. The ratings were using a 5-point Likert scale, where 1 indicates "very low" and 5 indicates "very high". To assess the knowledge level, the scores were averaged and categorized into five difficulty levels: very low, low, medium, high, and very high, as tabulated in Table 2 (Moidunny, 2009).

Table 2. Level of knowledge

Mean Score	Mean interpretation table
1.00 – 1.80	Very Low
1.81 – 2.60	Low
2.61 – 3.20	Medium
3.21 – 4.20	High
4.21 – 5.00	Very High

Source: Moidunny (2009)

With mean scores for all montage elements ranging from 3.9778 to 4.0519, the knowledge level is classified as "High," indicating the overall effectiveness of the hands-on video editing workshops. This suggests that the workshops were well-executed and served as an effective teaching approach. Furthermore, hands-on workshops provide a practical alternative to traditional educational methods, ensuring that students develop the essential skills needed to excel in video editing.

Table 3. Mean score for knowledge

Variables	Mean Score	Level of Difficulty
Visual	3.9778	High
Music	4.0000	High
Graphics	3.9852	High
Transition	4.0519	High
Rhythm	4.0296	High

4.3 Evaluate the relationship between knowledge and students' perceptions towards the hands-on video editing workshops

(i) Normality Test

A normality test is an essential analysis that is conducted to determine if the distribution of data is normally distributed or not, before proceeding to further analysis. Additionally, normality can be assessed using the Skewness value, where data with a Skewness value between -2 and +2 is generally considered normally distributed (Pallant, 2011).

Table 4. Summary of skewness value

Variable	Skewness
Knowledge	-0.369
Perception	-0.159

Table 4 summarizes the skewness values for two variables, namely knowledge and perception that are involved in this study. The result concluded that both variables are normally distributed, as their skewness values fall within the range of -2 to +2.

(ii) Reliability Analysis

The reliability of the item used in this study was assessed through a reliability test. According to Pallant (2011), an item is considered reliable if the Cronbach's Alpha value exceeds 0.7. As shown in Table 5, the Cronbach's Alpha values for both variables were above 0.7 (knowledge = 0.986 and perception = 0.979), confirming that the data used in this study is reliable.

Table 5. Reliability analysis

Variable	Number of Items	Cronbach's Alpha
Knowledge	5	0.986
Perception	7	0.979

(iii) Regression Analysis

A regression analysis was conducted to answer the second research objective. As tabulated in Table 6, the result reveals that knowledge of montage elements has a significant relationship with perception towards the Hands-On Video Editing Workshops (p -value < 0.05). This suggests that students' knowledge on montage elements can influence their perception towards the workshop.

Value of r^2 from Table 6 indicates that 30.8% of the variation in students' perception towards the workshops can be attributed to students' knowledge of montage elements. The remaining 69.2% is explained by other factors. Overall, the regression model shows significant results (p -value < 0.05). In summary, the significant relationship between knowledge of montage elements and participants' perceptions highlights the crucial role of montage techniques in influencing how participants assess and interact with video editing workshops.

Table 6. Regression Analysis

Variable	Result of Regression				
	Coefficient	t	P - value	r^2	F test
Constant	3.156	9.101	0.005	0.308	11.147 (0.003)
Knowledge	0.338	3.339	0.003		

5. CONCLUSION

At the end of this study, it was found that hands-on video-editing workshops significantly increase montage-element knowledge among non-technical students. The high mean scores suggest that students have found montage elements easy to learn, confirming the success of the workshops. Regression analysis showed a significant relationship between the familiarity of students with montage elements and their perception of the workshops; the more they knew, the better they perceived the experience. The study highlights the necessity of hands-on methods in video editing courses for an effective and interactive experience. However, this study is not without limitations. The relatively small sample size ($n=27$) and the short duration of the intervention may limit the generalizability of the findings. Future research should examine the long-term effects of hands-on workshops on skill retention and their applicability in real-world contexts. Moreover, studies may focus on fostering active learning through collaborative methods, the integration of advanced editing techniques, and the implementation of practical workshop sessions. Expanding the scope to include larger and more diverse student populations would further enhance the generalizability of the findings. Furthermore, integrating emerging technologies such as artificial intelligence into video editing education may further enhance student engagement and learning outcomes.

6. ACKNOWLEDGEMENTS/FUNDING

The authors would like to acknowledge the support of students who have participated in the video editing workshop for the March-October 2024 semester from the Diploma in Public Administration (AM110) and the Diploma in Halal Management (IC120) at Universiti Teknologi MARA Sarawak, Kampus Samarahan 2.

7. CONFLICT OF INTEREST STATEMENT

The authors declared that they have no conflicts of interest to disclose.

8. AUTHORS' CONTRIBUTIONS

Satria Arjuna Julaihi played a crucial role in shaping the project's direction by contributing to its initial conceptualization. **Norizuandi Ibrahim**, **Abdul Hadi Abdul Talip**, and **Zubaidah Bohari** were responsible for executing the workshop and contributed to the writing process. They collaborated on

<http://dx.doi.org/10.24191/jcrinn.v10i2.531>

drafting the original article and refining it through meticulous revisions. Meanwhile, **Rumaizah Che Md Nor** applied her expertise in data analysis, utilizing statistical and computational methods. Together, their combined efforts highlight a well-rounded and collaborative approach, showcasing the diverse skills and responsibilities each team member undertook to ensure the successful execution of the research.

9. REFERENCES

- Aggad, S. (2023). *Challenges, opportunities, and success factors of integrating video production education in Lebanon* [Doctoral dissertation, Lebanese American University]. SoAS – Theses and Dissertations. <https://doi.org/10.26756/th.2023.647>
- AlKhaibary, A. A., Ramadan, F. Z., Aboshaiqah, A. E., Baker, O. G., & AlZaatari, S. Z. (2020). Determining the effects of traditional learning approach and interactive learning activities on personal and professional factors among Saudi intern nurses. *Nursing Open*, 8, 327–332. <https://doi.org/10.1002/nop2.633>
- Anas, I. (2019). Behind the scene: Student-created video as a meaning-making process to promote student active learning. *Teaching English with Technology*, 19(4), 37–56.
- Brown, A., & Green, T. (2021). *The essentials of instructional design*. Routledge. <https://doi.org/10.4324/9780429439698>
- Dancyger, K. (2018). *The technique of film and video editing: History, theory, and practice* (6th ed.). Routledge. <https://doi.org/10.4324/9781315210698>
- DeGuzman, K. (2024, December 30). *What is a montage? Definition, examples & 6 ways to use them*. StudioBinder. <https://www.studiobinder.com/blog/what-is-a-montage-definition/>
- Dorfman, B. S., Terrill, B., Patterson, K., Yarden, A., & Blonder, R. (2019). Teachers personalize videos and animations of biochemical processes: Results from a professional development workshop. *Chemistry Education Research and Practice*, 20(4), 772–786. <https://doi.org/10.1039/C9RP00057G>
- Fleischmann, K. (2021). Hands-on versus virtual: Reshaping the design classroom with blended learning. *Arts and Humanities in Higher Education*, 20(1), 87–112. <https://doi.org/10.1177/1474022220906393>
- Galatsopoulou, F., Kenterelidou, C., Kotsakis, R., & Matsiola, M. (2022). Examining students' perceptions towards video-based and video-assisted active learning scenarios in journalism and communication courses. *Education Sciences*, 12(2), 74. <https://doi.org/10.3390/educsci12020074>
- Getenet, S., Cante, R., Redmond, P., & Albion, P. (2023). Students' digital technology attitude, literacy and self-efficacy and their effect on online learning engagement. *International Journal of Educational Technology in Higher Education*, 21, 3. <https://doi.org/10.1186/s41239-023-00437-y>
- Hsieh, C., & Knudson, D. (2020). Students' perception on teaching style and learning outcome. *ISBS Proceedings Archive*, 38(1), Article 7. <https://commons.nmu.edu/isbs/vol38/iss1/7/>
- Jeong, H. D. J., & Lim, J. (2024). Implementing a video production collaboration system with information marker overlay technique. *International Journal of Advanced Media and Communication*, 8(1), 77–92. <https://doi.org/10.1504/IJAMC.2024.140661>
- Kayyali, M. (2024). Digital literacy in higher education: Preparing students for the workforce of the future. *International Journal of Information Science and Computing*, 11(1), 53–73. <https://doi.org/10.30954/2348-7437.1.2024.6>

- Kirk, B. (2015). *A proactive, experiential and student-centered learning approach: A case study of the effects of a social media video editing "App" in a traditional classroom setting* [Doctoral dissertation, Liberty University]. Doctoral Dissertations and Projects. <https://digitalcommons.liberty.edu/doctoral/1084>
- Knoblauch, C. (2022, June). Experiential learning in digital contexts - A case study. In *The Learning Ideas Conference* (pp. 181–191). Springer. https://doi.org/10.1007/978-3-031-21569-8_17
- Kolb, D. A. (1984). *Experiential learning: Experience as the source of learning and development*. Prentice Hall.
- Merriam-Webster. (2025, February 22). *Montage*. Merriam-Webster.com dictionary. <https://www.merriam-webster.com/dictionary/montage>
- Moidunny, K. (2009). *The effectiveness of the national professional qualification for educational leaders (NPQEL)* [Unpublished doctoral dissertation, The National University of Malaysia].
- Nelson, G. L. (2018). *Perceptions about hands-on art making by online students* [Doctoral dissertation, Walden University]. ProQuest.
- Priandika, A. T., Permata, P., Gunawan, R. D., Ardiansah, T., Fahrizal, M., Maylani, A., & Anggraini, A. (2022). Video editing training to improve the quality of teaching and learning at SMK Palapa Bandarlampung. *J. Eng. Inf. Technol. Community Serv*, 1(2), 26–30. <https://doi.org/10.33365/jeit-cs.v1i2.134>
- Pallant, J. (2020). *SPSS survival manual: A step by step guide to data analysis using SPSS program*. Allen & Unwin. <https://doi.org/10.4324/9781003117452>
- Scheef, A. R., & Johnson, C. (2017). The power of the cloud: Google Forms for transition assessment. *Career Development and Transition for Exceptional Individuals*, 40(4), 250–255. <https://doi.org/10.48550/arXiv.1805.03522>
- Snelson, C. (2018). Video production in content-area pedagogy: A scoping study of the research literature. *Learning, Media and Technology*, 43(3), 294–306. <https://doi.org/10.1080/17439884.2018.1504788>
- So, J., & Chan, V. (2018). Impact of using online questionnaires in collecting student feedback in development activities. *Proceedings of the 2018 International Conference on Learning and Teaching in Computing and Engineering*, 31–34. <https://doi.org/10.1145/3183586.3183597>
- Sufi, S., Nenadic, A., Silva, R., Duckles, B., Simera, I., de Beyer, J. A., ... Higgins, V. (2018). Ten simple rules for measuring the impact of workshops. *PLoS Computational Biology*, 14(8), e1006191. <https://doi.org/10.1371/journal.pcbi.1006191>
- Theobald, E. J., Hill, M. J., Tran, E., Agrawal, S., Arroyo, E. N., Behling, S., ... Freeman, S. (2020). Active learning narrows achievement gaps for underrepresented students in undergraduate science, technology, engineering, and math. *Proceedings of the National Academy of Sciences*, 117(12), 6476–6483. <https://doi.org/10.1073/pnas.1916903117>
- Wan, X., Perumal, V., & Neo, T. K. (2022). A critical review on the use of montage technique in film and television. In *Proceedings of the 2nd International Conference on Creative Multimedia 2022* (pp. 189–196). Atlantis Press. https://doi.org/10.2991/978-2-494069-57-2_25
- Xiang, W., Perumal, V., & Neo, T. K. (2023, February). A critical review on the use of montage technique in film and television. In *Proceedings of the 2nd International Conference on Creative Multimedia 2022*. https://doi.org/10.2991/978-2-494069-57-2_25

- Yen, C. S., & Yang, S. C. (2024). The integration of artificial intelligence and video production skills in workplace development: A study from the perspective of vocational training. *The Review of Socionetwork Strategies*, 18(2), 373–386. <https://doi.org/10.1007/s12626-024-00162-6>
- Zhang, I. Y., Tucker, M. C., & Stigler, J. W. (2022). Watching a hands-on activity improves students' understanding of randomness. *Computers & Education*, 186, 104545. <https://doi.org/10.1016/j.compedu.2022.104545>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Analysing the Potential Vulnerabilities in Online Voting Protocols: Homomorphic Encryption Approach

Yao Charles Azameti^{1*}, Wiiliam Asiedu², George Asante³

^{1,2,3}Akenten Appiah-Menka University of Skills Training and Entrepreneurial Development, Kumasi, Ghana.

ARTICLE INFO

Article history:

Received 14 March 2025

Revised 21 April 2025

Accepted 23 April 2025

Published 1 September 2025

Keywords:

Blockchain Voting

Homomorphic Encryption

Electronic Voting Security

DOI:

10.24191/jerinn.v10i2.515

ABSTRACT

Online voting offers a hopeful alternative to traditional voting approaches by improving accessibility, efficiency, and transparency in elections. However, despite their potential, these systems are prone to a variety of security vulnerabilities. This research aims to analyze the crucial vulnerabilities that can compromise the integrity, confidentiality, and availability of online voting systems. Key threats examined include cyber-attacks such as denial of service (DoS), man-in-the-middle (MITM) attacks, and malware injection, as well as issues related to voter authentication, anonymity, and coercion resistance. Additionally, the research evaluates the effectiveness of cryptographic techniques like homomorphic encryption, zero-knowledge proofs, and blockchain in mitigating these risks. The study provides a comprehensive review of existing protocols, identifies gaps in their security architectures, and proposes enhanced mechanisms to address the vulnerabilities. The ultimate goal is to contribute to the development of more robust and secure online voting systems that ensure voter trust and uphold democratic principles.

1. INTRODUCTION

Casting votes online has emerged as a viable alternative to traditional voting systems, offering enhanced convenience, efficiency, and transparency. As societies move towards digitalization, government and organizations alike are increasingly exploring the potential of online voting systems to conduct elections. These systems promise to reduce logistical barriers, improve voter turnout, and provide immediate results, all while minimizing the costs associated with physical polling stations and paper ballots. Despite these advantages, the widespread adoption of online voting has been slow, primarily due to concerns about the security and reliability of these systems. The significance of online voting lies in its ability to maintain integrity, transparency, and security within the electoral process. By integrating modern digital technologies, online voting protocols aim to ensure the accuracy of electoral data, minimize the risk of tampering or fraud, and enhance public trust in election outcomes.

This section will delve into research works in the field of online voting protocols, examining the evolution of design paradigms, encryption techniques, and implementation strategies. A key challenge of online voting protocols is their reliance on a central point of control for vote casting and validation. Unlike decentralized systems, where multiple nodes independently verify transactions, most online voting systems are centralized, with votes processed through a single authoritative entity. This centralization usually creates vulnerabilities, as it becomes a single point of failure that could be targeted by cyberattacks or internal manipulation. Additionally, the trust placed in a single authority can reduce transparency and make it harder for participants to independently verify the integrity of the election process. This reliance on centralized systems raises concerns about the resilience of online voting to malicious attacks and

^{1*} Corresponding author. E-mail address: ytornyeviadi@gmail.com

undermines the potential for trust among participants. Moreover, cryptographic techniques play a pivotal role in ensuring the security and privacy of votes in online voting protocols. Advanced cryptographic methods such as zero-knowledge proofs, homomorphic encryption, and multi-party computation allow voters to cast their ballots anonymously while ensuring the verifiable aggregation of votes. By integrating these cryptographic tools into the voting process, online voting protocols aim to strike a balance between privacy and transparency, thereby safeguarding the integrity of elections.

2. LITERATURE SURVEY

The following review of literature provides a comprehensive overview of the evolution of online voting systems, from their conceptual foundations to practical implementation. By examining the strengths, limitations, and challenges of existing systems, insights will be gleaned into best practices and areas for improvement in the development of an efficient and secure voting systems. This review will inform the subsequent chapters, building upon prior research to propose a novel voting protocol that addresses the complexities and exigencies of modern democratic processes.

2.1 Online voting system

Online Voting, also referred to as E-Voting is a method of casting votes using electronic devices connected to the internet, typically allowing voters to participate in elections remotely. The system facilitates greater accessibility by enabling voting from any location with an internet connection, reducing the need for physical polling stations. Online voting protocols often integrates various security features, such as encryption, digital signatures, and biometric verification, to ensure the integrity, confidentiality, and authenticity of the voting process. E-Voting promised to increase voter turnout and efficiency, online voting also faces significant challenges, including security risks, digital divide issues, and concerns about voter privacy and coercion. The shift from traditional paper-based voting methods to internet-based voting (I-voting) has gained significant attention in recent years due to its potential to improve electoral processes.

Smith and Clark's research delves into the transition from traditional voting methods to internet-based voting (I-voting), underscoring benefits like cost reduction, improved vote accuracy, and greater accessibility, while also addressing significant challenges such as the digital divide, which limits access to online voting for certain demographics, and security risks including potential cyberattacks, voter fraud, and data breaches; they advocate for the development of standardized electronic voting (e-voting) systems to mitigate these risks but note that barriers such as inadequate security protocols, the complexity of voter authentication, and lack of public trust pose significant obstacles to widespread adoption of I-voting. Despite these challenges, they argue that I-voting could enhance voter participation and democratization, provided that these issues are carefully addressed (Smith & Clark, 2005). Thakur et al. (2014) explored the evolution of voting systems, highlighting the transition from traditional methods to mobile voting (m-voting) using Near Field Communication (NFC) and biometric verification. Their study addresses challenges like low voter turnout, particularly among youth, and proposes a mobile voting model that leverages common-off-the-shelf (COTS) mobile phones. Their proposed model enhances voter mobility, transparency, and ease of use while mitigating security risks. This approach aims to increase electoral engagement and accessibility across diverse demographics, making voting more efficient and inclusive (Thakur et al., 2014). In another research conducted by Al-Shammari et al, discusses the advantages of electronic voting systems (E-voting), highlighting their potential to outperform traditional voting methods through enhanced accessibility, reduction of voter errors, and cost-effectiveness. It emphasizes how features such as audio interfaces for visually impaired voters and simplified ballot management contribute to a more efficient voting process. Additionally, it points out the significant savings in costs associated with the use of DRE voting machines compared to paper ballots (Al-Shammari et al., 2012).

Another paper presented by Ahmad et al. provide a comprehensive overview of the evolution of electronic voting systems, categorizing them into four main types: punch card, optical scanning, direct recording electronic (DRE), and remote internet voting, while highlighting the benefits such as improved transparency, faster processing, enhanced security, and greater accessibility for remote and disabled voters. However, they identify several vulnerabilities and challenges, including susceptibility to cyberattacks,

tampering with digital infrastructure, privacy concerns with remote internet voting, and potential software malfunctions in DRE systems. The authors stress that while e-voting systems offer significant advantages, public trust can only be maintained through addressing these security risks, ensuring the integrity of election data, and implementing stronger verification mechanisms to prevent manipulation or fraud (Ahmad et al., 2021). A Survey conducted by Mursi et al. (2013), provided a comprehensive review of the evolution and state of electronic voting systems. The paper outlines the challenges associated with transitioning from traditional paper-based voting methods to electronic systems, focusing on the conflicting requirements of privacy, security, transparency, and fairness. The survey categorizes various e-voting schemes, discusses their cryptographic underpinnings, and examines the advantages and disadvantages of different approaches.

Key security concerns include voter authentication, vote integrity, privacy, and resistance to coercion. The paper further details the vulnerabilities of both traditional and modern voting systems, including punch cards, optical scanners, DRE (Direct Recording Electronic) systems, and online voting. The authors discuss the role of cryptographic mechanisms like homomorphic encryption, zero-knowledge proofs, and visual cryptography in enhancing the security of electronic voting. The survey concluded with a comparative analysis of different e-voting schemes, highlighting gaps in current technologies, especially in terms of scalability and trustworthiness and advocates for improvements in cryptographic protocols and suggests that biometric tokens and robust end-to-end verifiability are promising avenues for future secure e-voting systems (Mursi et al., 2013). In the quest to modernize India's voting process, Ganesh Prabhu et al, presents a new method also known as Smart Online Voting System to allow citizens to vote remotely via an online platform that uses facial recognition and OTP (One-Time Password) authentication.

This system eliminates the need for physical presence at polling stations, enabling users to vote from anywhere using a computer or mobile phone. It also offers an offline option where voters can use RFID tags instead of traditional voter IDs. The voting process involves two-step authentication, first through facial recognition and then by verifying an OTP sent to the registered mobile number. Results are updated in real-time in a central database, ensuring quick access and minimizing vote tampering. Despite its innovations, the system has several issues as it relies heavily on stable internet and technology, which may not be available to all. The system as well faces security risks like cyber-attacks, phishing, and denial-of-service (DoS) attacks. Again, storing biometric data in a central database raises privacy concerns that, the central database could be the point of failure or attack, and also the potential for voter coercion or fraud, especially in remote voting situations, is a significant challenge.

Thus, while the system enhances convenience and efficiency, addressing these security and privacy concerns became critical to its success (Ganesh Prabhu et al., 2021). To address the persistent challenges associated with online voting systems, such as security vulnerabilities, voter privacy concerns, accessibility issues, and lack of transparency, another paper proposes a more advanced solution that integrates enhanced security protocols and innovative technologies aimed at improving the security, accessibility, and convenience of elections by allowing citizens to vote online from any location. This system is designed to ensure high security using a combination of cryptographic techniques, face recognition, and password protection, which together form a robust authentication process. Users log in with a unique voter ID generated by the Electoral Commission of India, and their votes are cast through a web interface after verifying their identity with both a password and a facial image stored in the database. The system aims to increase voter turnout by making the process more accessible, especially for citizens in remote locations or those with mobility issues. The online voting system promises to speed up the counting process, reduce human errors, and eliminate vote tampering or rigging, offering a more transparent and reliable method of conducting elections.

The integration of face recognition technology further strengthens the security of the system by ensuring that only verified users can vote, yet it faces several crucifixions since the system remains vulnerable to cyberattacks such as hacking or denial-of-service (DoS) attacks that poses privacy concerns due to the storage of sensitive data like facial images in a central database, also the system has been criticized of being difficult for non-tech-savvy users or those without reliable internet access to use. Additionally, voters must place their trust in the technology, as the process lacks transparency, which led to concerns about the accuracy and security of their vote compared to traditional methods (Kaliyamurthi et al., 2013).

2.2 The need for blockchain

In online voting systems, centralization has long been a point of vulnerability. Centralized authorities whether election officials, government bodies, or electronic voting machines hold significant control over the voting process, from voter registration to vote tallying and result verification. This centralization creates opportunities for manipulation, fraud, and even cyber-attacks, as seen in numerous instances of compromised electronic voting machines and tampered elections.

The reliance on a central authority and database also introduces trust issues, as citizens must trust that officials and their systems are secure, impartial, and functioning properly. At the heart of it all is blockchain technology, which promised to offer a decentralized alternative, which addresses these critical concerns. By distributing the control of the voting ledger across multiple nodes, blockchain removes the reliance on any single authority. Each transaction, or vote, is encrypted, time-stamped, and verified by a network of nodes, ensuring that no single party can alter the results without detection. This decentralized structure not only provides transparency and security but also ensures the integrity of the voting process. Blockchain's tamper-resistant ledger, public verifiability, and resilience to attacks make it an ideal solution for transforming how elections are conducted, offering a modern, scalable solution to the vulnerabilities of centralized voting systems.

Osgood (2016), explores the need for blockchain in voting systems, highlighting the current vulnerabilities in electronic and paper-based voting, such as susceptibility to fraud, hacking, and lack of scalability, and argues that blockchain offers a secure, tamper-proof, and transparent alternative by distributing voting data across a decentralized ledger. In this paper, Osgood point out few weaknesses in the proposed blockchain voting protocols which include the challenge of securing voter authentication, risks of voter intimidation in remote voting, the potential for network attacks, and the reliance on internet infrastructure, which may limit the system's feasibility in certain regions (Osgood, 2016).

Again, a blockchain-based e-voting system that aims to improve the security, transparency, and privacy of elections by utilizing a private, permissioned blockchain infrastructure has been proposed by (Hjalmarsson et al., 2018). This system benefits from key blockchain strengths such as immutability, verifiability, and decentralized consensus, ensuring that votes are tamper-proof and publicly auditable, yet anonymous, as each vote is appended to a distributed ledger only after consensus is reached by multiple nodes. It also introduces smart contracts to automate election processes, including voter registration, vote tallying, and verification, reducing human error and central authority involvement, which lowers the potential for fraud and manipulation.

However, the system is not without critics. While the blockchain provides transparency, it does not inherently ensure voter authentication, requiring additional government identity verification services, which introduces points of vulnerability. Again, the system is challenged in resisting voter coercion, particularly in remote or unsupervised voting environments, limiting its suitability for large-scale, unsupervised elections. Also, the reliance on blockchain technology brings concerns about scalability, as high transaction throughput may be required for national elections, and existing blockchain implementations may struggle with performance under high voter turnout. Additionally, the proposed system faces the potential risk of a 51% attack, where a malicious entity could gain control of the majority of network nodes to manipulate results. This vulnerability is a known issue in public blockchains. However, the paper addresses this concern by employing a permissioned blockchain, which restricts node participation to trusted institutions, thereby reducing the likelihood of such an attack.

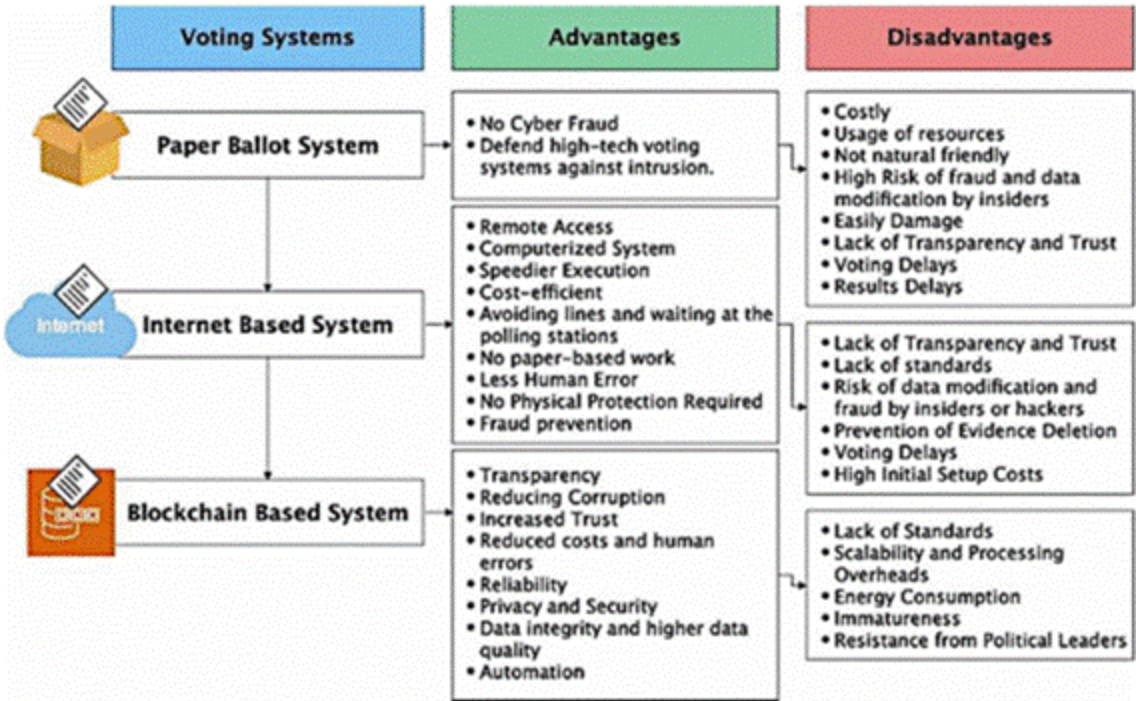


Fig. 1. Blockchain: The blockchain structure

Source: Jafar et al. (2022)

Fig. 1 compares three voting systems: Paper Ballot, Internet-Based (Online), and as well Blockchain-Based, concluding on the fact that while paper ballots are secure from cyber fraud but costly and inefficient, internet-based voting systems offer remote access and efficiency but face transparency and security issues, and blockchain-based voting systems provide enhanced transparency, security, and automation but encounter challenges with scalability, energy consumption, and high setup costs. The research further proposes a cost-efficient and scalable e-voting system based on Ethereum blockchain to address the shortcomings of conventional voting methods, such as lack of transparency, security, and scalability. The paper argued that blockchain can ensure immutable, tamper-proof of election data while reducing costs and improving efficiency, making it suitable for large-scale elections by employing off-chain solutions and sharding techniques to handle high volumes of transactions. However, the protocol faces several issues, including reliance on trusted authorities to manage the off-chain components, potential security risks from quantum attacks, and related issues with scalability as the blockchain grows, which may turn to increase latency and storage requirements (Jafar et al., 2022).

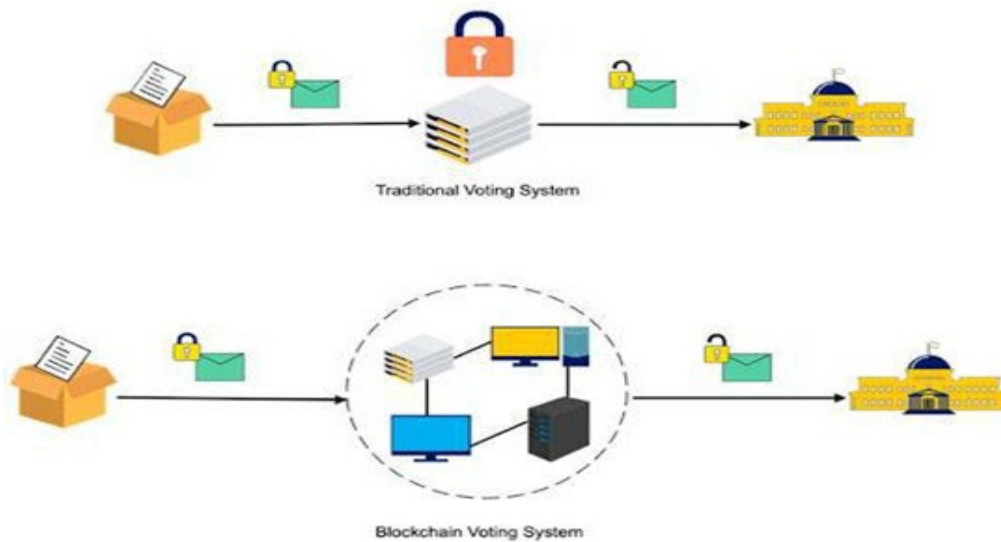


Fig. 2. Traditional (e-voting) vs. blockchain voting system

Source: Jafar et al. (2021)

The aim of using such a data structure is to achieve provable immutability of the blockchain. If any piece of data is altered, the block's hash containing this piece needs to be recalculated or be updated, and the hashes of all subsequent blocks also need to be recalculated (Nofer et al., 2017). This means only the hash of the latest block has to be used to guarantee that all the data remains unaltered. In blockchain solutions, data stored in blocks are formed from all the validated transactions during their creation, in this regards no one can insert, delete or alter the transactions within a block that has already undergone validation without it being noticed (Stephen & Alex, 2018). The initial zero-block, called the “genesis block,” usually contains some network settings, for example, the initial set of validators (those who issue the blocks). In further reviews, blockchain technology addressed flaws in the current electoral process by making the polling process transparent and easily available, preventing fraudulent voting, bolstering data security, and verifying poll results. The centralized nature of blockchain prevent tampering or manipulation at the server level, blockchain systems inherently protect against such risks by recording votes on a distributed, immutable ledger. Even though Helios emphasizes individual and universal verifiability, its reliance on a single server contrast with the decentralized and trustless nature of blockchain, which eliminates the need for users to the voting processes (Xiao et al., 2020).

The primary distinction between the two voting systems is seen in Fig. 2. In the traditional online voting systems, votes are cast through a central authority, which makes it easy for someone to alter or change a record; nobody is aware of how to verify that record. In contrast, in decentralized voting systems, data is stored across multiple nodes, making it impossible for someone to manipulate every node and alter the data. As a result, votes cannot be destroyed and can be effectively verified by tallying with other nodes. Elections are won at the polling stations, according to electoral regulations, hence Agbesi and Asante (2019), concluded that if the system permits figures to be manipulated, it could tarnish the reputation of the election process. The problem of manipulating polling station results has led to mistrust in the electoral body and unrest in a number of African countries. However, this research believes that the blockchain revolution presents an opportunity to include auditability, integrity, trust, and transparency into the transmission and storing of results from many locations. These specifications can be met since blockchain technology allows for the creatin of decentralized systems, and multiple stakeholders will own the database rather than just one system administrator in the case of traditional e-voting (Kshetri & Voas, 2018). A detailed literature reviewed of Helios Online voting protocol proposed by Adida (2008), revealed that Helios voting protocol, while offering robust security and transparency through cryptographic techniques such as homomorphic encryption, the ElGamal cryptosystem, and zero-knowledge proofs, relies on a

centralized server for tallying votes, which paved way for potential vulnerabilities related to trust and single points of failure. Instead of implementing blockchain to distribute votes across multiple nodes in a decentralized network, ensuring immutability and transparency through an open ledger, Helios remains dependent on a central authority for vote counting and validation.

2.3 Homomorphic encryption

Based on the history of cryptology, homomorphism was first put out as a possible remedy for the computation without decrypting problem in blockchain voting systems by Rivest et al. (1978) back in 1978. Many attempts were made by analysts worldwide to develop a homomorphic scheme with few operations after Rivest et al. A homomorphic encryption method offers a way to calculate encrypted data directly while maintaining privacy. Homomorphic encryption was used to conduct calculations on ciphertext, producing an encrypted output in the process (Probor et al., 2023). In order to perform computations on data, the majority of encryption systems permit third parties to decrypt the data. Anytime a third party is involved in data computation, security issues arise. It is greatly desired to have an encryption scheme that permits calculation on encrypted data without decryption (Al Badawi et al., 2021). These kinds of encryptions systems are known as homomorphic encryptions. When calculations are performed on encrypted data, homomorphic encryption ensures that the results are the same to those obtained from calculations performed on unencrypted data (Tebaa et al., 2012). A previous definition of homomorphic encryption was an encryption technique in which an algebraic operation on the plaintext is equal to another algebraic operation on the ciphertext (Ogburn et al., 2013). Wu and Haven carried out two experiments in Using variance and mean of massive sets of encrypted data through linear regressions. The number of Homomorphic Encryption for massive Scale Statistical Analysis, wherein they computed the data points in these tests rose to one million and four million elements, respectively (Wu, 2012).

2.4 Homomorphic encryptions in elections

Recently, Homomorphic encryption has been used in the development of online voting platforms and this is driven by the requirement to create cutting-edge security measures in order to enable the widespread adoption of online voting around the globe (Alvarez & Hall, 2003). In the year 2010, homomorphic encryption-based voting was introduced by George and Sebastian, the plan succeeds in achieving receipt-freeness, secrecy, and coercibility. Both yes/no and multi-candidate voting can be conducted using this approach (George & Sebastian, 2010). Huszti presented a voting mechanism based on homomorphic encryption and the Cramer mechanism. Andrea Huszti's voting mechanism integrates homomorphic encryption with the Cramer mechanism, creating a secure electronic voting system that allows encrypted votes to be tallied without ever needing decryption, which intended to maintain voter confidentiality and data integrity. The system achieves essential properties such as eligibility (ensuring only authorized voters can vote), privacy (votes cannot be traced back to individuals), verifiability (voters can confirm their vote was counted correctly without revealing it), receipt-freeness (preventing voters from proving how they voted to avoid coercion), and coercibility resistance (ensuring voters cannot be forced to vote a certain way). Votes are encrypted at the point of casting, and homomorphic encryption allows vote aggregation directly on encrypted data, preserving the privacy of individual votes. The use of anonymous channels further ensures that voters' identities remain hidden throughout the process.

However, despite its robust cryptographic underpinnings, the system is vulnerable to significant security risks. First, side-channel attacks techniques that exploit physical characteristics of the system, such as timing information, power consumption, or electromagnetic leaks could potentially extract sensitive information during vote encryption or tallying. Also, the security of the system hinges on the assumption that authorities managing the decryption keys are trustworthy; any collusion among these authorities or compromise of key management protocols could result in vote manipulation or exposure of voter identities. Additionally, key leakage, whether through insider threats, poor key management, or advanced cryptographic attacks, could undermine the encryption process's security, allowing adversaries to decrypt votes. In all, the reliance on anonymous channels poses a challenge; if these channels are compromised or improperly implemented, it could lead to the de-anonymization of voters, breaking the privacy and receipt-freeness guarantees of the system (Huszti, 2011). A novel voting technique based on the blind signature RSA and the additive homomorphic characteristic of the Paillier cryptosystem was proposed by (Hussien

& Aboelnaga, 2013), effectively achieves eligibility, confidentiality, privacy, uniqueness, and correctness. However, the technique remains vulnerable to potential attacks on key management, possible collusion between authorities compromising voter anonymity, and susceptibility to cryptographic weaknesses in the Paillier cryptosystem, such as chosen ciphertext attacks, which could undermine the security guarantees if not properly mitigated.

Similarly, Yi and Okamoto (2013) proposed a voting technique that ensures voter privacy even under physical coercion or malware attacks on the voter's device, by only revealing the election outcome (win or lose) without disclosing the specific count of yes or no votes; however, several attacks including the potential for side-channel attacks, where adversaries could infer voting patterns through indirect data leaks, and susceptibility to software and hardware manipulation that could compromise the integrity of the results or allow tampering without detection.

A voting system based on homomorphic encryption was presented at a conference session that was captured by Zhao (2014) in order to guarantee anonymity, privacy, and dependability by using homomorphic encryption with the RSA cryptosystem to securely encrypt voting data; however, the system is vulnerable to certain attacks, such as chosen-ciphertext attacks (CCA), key exposure risks, and the inefficiency of RSA in handling large datasets, which could affect scalability and performance, and it also relies on the trustworthiness of key distribution and management. A partly homomorphic cloud-based mobile voting system was described in another conference proceedings published by Will et al. (2015). To demonstrate the system's usefulness, the paper put it into practice. Qualification, non-reuse, non-traceability, verifiability, accuracy of tally, non-coercibility, auditability, accessibility, equity, soundness, and integrity are all attained by the system. However, the system remains susceptible to threats like potential insider attacks, cloud infrastructure breaches, denial-of-service attacks, weaknesses in cryptographic protocols, and privacy risks during data transmission, which could compromise voter anonymity, the integrity of the voting process, and overall system trustworthiness. According to Yang et al. (2018), each ballot is encrypted using the exponential ElGamal cryptosystem before submission in order to safeguard the confidentiality of the votes. Additionally, the system makes sure that proofs are created and kept for every vote element during the voting process. Before counting, these proofs can be used to confirm each ballot's eligibility and validity without having to decode or read the ballot's content. However, their protocol is never without risks which include susceptibility to chosen-ciphertext attacks, potential inefficiency with large datasets due to RSA's computational overhead, and the reliance on RSA's key length, which could be weakened by advances in quantum computing.

Ravindran and Kalpana (2013) also argued that more capacity is needed for both encryption and re-encryption. In order to pack the encrypted data, data packing is utilized. They proposed a secure voting platform where votes are encrypted, re-encrypted, packed (zipped), and transmitted over insecure channels to ensure privacy, fairness, robustness, and verifiability. Upon receiving the encrypted votes, the system validates each vote, and if valid, proceeds with decryption and unpacking to tally the result and determine the winner. The voter's credentials are verified by the platform's verifier authority, ensuring eligibility compliance while maintaining individual and universal verifiability. Despite these security features, the process could be vulnerable during the packing and unpacking processes, as improper implementation could allow data leakage or manipulation. Furthermore, transmitting encrypted votes over insecure channels may expose the encrypted votes to man-in-the-middle attacks or replay attacks. And if the verifier authority is compromised, it could lead to unauthorized access or tampering with voter credentials or votes, undermining the system's security guarantees. Balasubramanian and Jayanthi (2016) introduced a secure voting system that employs the Paillier cryptosystem, known for its additive homomorphic encryption, allowing the tallying of encrypted votes without revealing individual vote content. Voters are assured that their vote is recorded in the Voting Table, enabling them to verify whether their vote was counted while maintaining the confidentiality of each vote. The system restricts participation to eligible voters, ensuring compliance with election requirements.

Despite these features, vulnerabilities exist, particularly with the integrity of the Voting Table, as it serves as the central point of vote verification and could be susceptible to tampering, fraud, or unauthorized access, potentially undermining the accuracy of the vote count. Additionally, the homomorphic encryption used, while secure, can introduce risks such as partial data leakage, where attackers could exploit patterns in the encrypted data or perform side-channel attacks to infer sensitive information about the votes.

Moreover, if any component in the cryptographic process is flawed or compromised, such as improper key management or weak encryption parameters, the confidentiality and correctness of the entire voting system could be jeopardized.

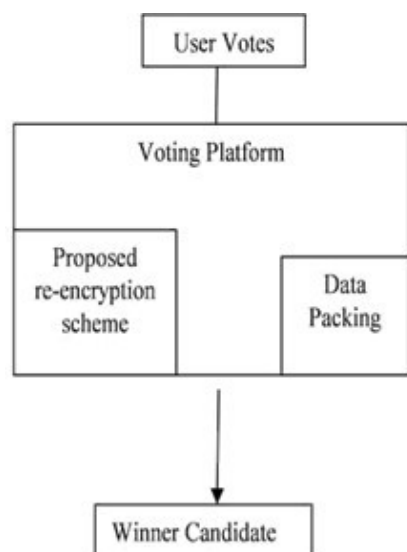


Fig. 3. Flow chart of the proposed system

Source: Balasubramanian and Jayanthi (2016)

2.5 Integration of homomorphic encryption and blockchain in voting protocol

In 1978, Blakley (1979) independently introduced methods for dividing a secret into multiple shares that can be distributed among mutually distrustful agents. This paper outlines a homomorphic property achieved by these and other secret sharing schemes, enabling the combination of multiple secrets through direct computations on the shares. This property minimizes the need for trust among agents and expands the applicability of secret sharing to various new challenges. One application presented here offers a simpler and more efficient approach to verifiable secret sharing compared to earlier methods. Another application described provides a fault-tolerant solution for conducting verifiable secret-ballot elections (Benaloh, 1987).

Jabbar and Alsaad (2017) introduced their first voting techniques based on homomorphic encryption, which allows mathematical operations to be performed directly on encrypted votes (ciphertext) without decrypting them, enabling privacy-preserving tallying in which the votes remain confidential throughout the process; however, with this protocol the possibility of side-channel attacks during the cryptographic operations were later observed, where an attacker could infer information about the encrypted data, as well as challenges related to key management, such as improper handling or exposure of encryption keys that could compromise the system's security. Additionally, if the homomorphic encryption scheme is not properly implemented or the cryptographic parameters are weak, it may result in inefficiencies, performance bottlenecks, or weakened resistance to cryptographic attacks, which could be exploited to manipulate or disrupt the election process.

Qu et al. (2022) proposed a blockchain-based electronic voting system utilizing homomorphic signcryption, which eliminates the need for a traditional trusted third party by using smart contracts to publicly manage the voting process on the blockchain, with ballots encrypted and signed through homomorphic encryption and signcryption algorithms. These algorithms allow the system to efficiently

aggregate votes and produce a homomorphic tally, improving voting efficiency and reducing the computational burden on voters, making the system scalable and practical for large-scale elections. However, there are potential risks associated with smart contract bugs or exploits, which could lead to unauthorized access, tampering, or even denial-of-service attacks on the voting process. Additionally, although blockchain offers transparency, if the encryption keys or signcryption processes are compromised, the confidentiality of votes could be at risk, leading to potential vote manipulation or privacy breaches. And also, relying on blockchain's consensus mechanisms alone might expose the system to vulnerabilities such as 51% attacks, which could disrupt the integrity of the voting process or prevent accurate tallying of results, hence the need to combine the two strong techniques to protect the integrity of each vote cast (Sayeed & Marco-Gisbert, 2019) .

Fan et al. (2020) proposed a voting system where voters use digital signature algorithms to sign their ballots before submission, ensuring that each voter can verify their own vote, and the system can authenticate the origin of the vote; however, as the number of voters and candidates increases, the computational burden and complexity of the ballot tallying and verification process also grow significantly, potentially leading to performance bottlenecks. The risk of signature forgery if the digital signature scheme is not robust enough may occur, the possibility of denial-of-service attacks targeting the computational resources needed for verification, and the increased exposure to replay or man-in-the-middle attacks during the transmission of signed ballots, especially if secure channels are not properly enforced. If the signature verification process is not efficiently scaled, the system may become vulnerable to delays or inaccuracies in tallying, compromising the timeliness and reliability of the election outcome.

3. RESULT

Upon numerous reviews of the existing works in online voting protocols, the findings reveal that existing online voting protocols face significant vulnerabilities, which include centralization risks that create single points of failure, a lack of transparency that undermines voter trust, susceptibility to security breaches such as side-channel and chosen-ciphertext attacks, and scalability challenges that hinder their performance in large-scale elections. Security weaknesses such as susceptibility to Denial of Service (DoS) attacks, malware injections, and Man-in-the-Middle (MITM) attacks were identified, particularly in centralized systems where a breach could lead to widespread manipulation. Privacy concerns were evident, as many protocols failed to guarantee voter anonymity and confidentiality, exposing voters to coercion and vote-selling. Additionally, many systems lacked verifiability and transparency, making it difficult for voters to confirm that their votes were accurately counted, which erodes trust in the system. Scalability was another issue, with existing systems struggling to handle large-scale elections efficiently, leading to delays and security risks.

Furthermore, most systems failed to implement sufficient coercion resistance mechanisms, leaving voters vulnerable in unsupervised settings. Weak voter authentication also opened the door for identity theft and impersonation. Blockchain technology, while addressing some of these issues by decentralizing voting processes and improving transparency, still faced scalability challenges and risks such as a 51% attack. Although blockchain increases transparency, it does not inherently solve voter anonymity, authentication, or coercion resistance issues. Similarly, homomorphic encryption, which allows encrypted data to be processed without decryption, has shown promise in securing vote tallying, but its high computational complexity presents challenges in large-scale elections. In all, in as much as online voting systems have improved accessibility and efficiency, they still suffer from significant vulnerabilities, especially in the areas of security, privacy, scalability, and voter authentication.

Technologies like blockchain and homomorphic encryption offer potential solutions, but they require further development to overcome challenges related to performance and scalability. The findings emphasize the need for a new approach that integrates enhanced security measures, such as robust voter authentication, decentralized consensus mechanisms, and scalable cryptographic protocols. There is the need to proposed

a hybrid blockchain-based voting protocol with homomorphic encryption aims to address these issues and create a more secure, transparent, and trustworthy system for large-scale elections.

4. CONCLUSION

The analysis conducted throughout this study highlights the persistent vulnerabilities and technical limitations of existing online voting systems, particularly regarding security, scalability, privacy, and user trust. Despite notable advancements in cryptographic techniques and system design, issues such as centralized control, susceptibility to cyber threats, and lack of verifiability continue to impede the full realization of secure and trustworthy online elections. While blockchain and homomorphic encryption each offer promising solutions, they also come with their respective challenges: blockchain with its scalability and privacy limitations, and homomorphic encryption with its computational intensity. To address these gaps, future research should focus on the development and refinement of hybrid voting protocols that integrate the strengths of both technologies: blockchain for its decentralized and tamper-resistant infrastructure, and homomorphic encryption for privacy-preserving vote tallying. Such hybrid solutions should aim to minimize computational overhead while ensuring real-time scalability, making them viable for national and large-scale elections. Emphasis should also be placed on addressing known cryptographic risks such as side-channel attacks, key leakage, and collusion among verification authorities. Moreover, beyond the technical domain, there is a pressing need for interdisciplinary studies that examine the social, legal, and operational challenges surrounding the deployment of secure online voting systems. These include issues related to voter accessibility, digital literacy, inclusivity, and public perception. Systems must be designed with a user-centric approach to foster trust, ensuring transparency and verifiability without compromising usability or voter anonymity. In conclusion, the transition to secure, scalable, and transparent online voting systems is both a technical and societal endeavor. A multidimensional approach that combines cutting-edge cryptography, robust system architecture, and thoughtful policy design is essential. By advancing in these directions, future voting systems can uphold democratic integrity and adapt to the evolving landscape of digital governance.

5. ACKNOWLEDGEMENTS/FUNDING

The authors would like to thank AAMUSTED community, FASME, and all lecturers of the IT Department for their support during the research. The authors did not receive any specific funding for this work.

6. CONFLICT OF INTEREST STATEMENT

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

7. AUTHORS' CONTRIBUTIONS

The authors were responsible for all aspects of the research including conceptualization, methodology, data collection, analysis, and writing of the manuscript.

8. REFERENCES

- Adida, B. (2008). Helios: Web-based open-audit voting. In *USENIX Security Symposium* (Vol. 17, pp. 335–348).
- Agbesi, S., & Asante, G. (2019). Electronic voting recording system based on blockchain technology. In *2019 12th CMI Conference on Cybersecurity and Privacy (CMI)* (pp. 1–8). IEEE. <https://doi.org/10.1109/CMI48017.2019.8962142>

- Ahmad, A. H., Tabassum, F. N., Ayaz, S. A., Bahir, N., & Meelam, M. (2021). Electronic voting system: Nature, origin and its global application. *International Journal of Innovation, Creativity and Change*, 15(2), 1334–1337.
- Al Badawi, A., Polyakov, Y., Aung, K. M. M., Veeravalli, B., & Rohloff, K. (2021). Implementation and performance evaluation of RNS variants of the BFV homomorphic encryption scheme. *IEEE Transactions on Emerging Topics in Computing*, 9(2), 941–956. <https://doi.org/10.1109/TETC.2019.2902799>
- Al-Shammari, A. F. N., Villafiorita, A., & Weldemariam, K. (2012). Understanding the development trends of electronic voting systems. In *2012 Seventh International Conference on Availability, Reliability and Security* (pp. 186–195). IEEE. <https://doi.org/10.1109/ARES.2012.76>
- Alvarez, R. M., & Hall, T. E. (2003). *Point, click, and vote: The future of Internet voting*. Rowman & Littlefield.
- Balasubramanian, K., & Jayanthi, M. (2016). A homomorphic crypto system for electronic election schemes. *Circuits and Systems*, 07(10), 3193–3203. <https://doi.org/10.4236/cs.2016.710272>
- Benaloh, J. C. (1987). Secret sharing homomorphisms: keeping shares of a secret (Extended Abstract). In A. M. Odlyzko (Ed.), *Advances in Cryptology --- CRYPTO' 86* (pp. 251–260). Springer Berlin Heidelberg.
- Blakley, G. R. (1979). Safeguarding cryptographic keys. *Managing Requirements Knowledge, International Workshop On* (pp. 313–313). IEEE Computer Society. <https://doi.org/10.1109/MARK.1979.8817296>
- Fan, X., Wu, T., Zheng, Q., Chen, Y., Alam, M., & Xiao, X. (2020). HSE-Voting: A secure high-efficiency electronic voting scheme based on homomorphic signcryption. *Future Generation Computer Systems*, 111, 754–762. <https://doi.org/10.1016/j.future.2019.10.016>
- Ganesh Prabhu, S., Nizarahammed., A., Prabu., S., Raghu., S., Thirrunavukkarasu, R. R., & Jayarajan, P. (2021). Smart online voting system. In *2021 7th International Conference on Advanced Computing and Communication Systems (ICACCS)* (pp. 632–634). IEEE. <https://doi.org/10.1109/ICACCS51430.2021.9441818>
- George, V., & Sebastian, M. (2010). An adaptive indexed binary search tree for efficient homomorphic coercion resistant voting scheme. *International Journal of Managing Information Technology*, 2(1), 1–9.
- Hjalmarsson, F. P., Hreioarsson, G. K., Hamdaqa, M., & Hjalmytsson, G. (2018). Blockchain-based e-voting system. In *IEEE International Conference on Cloud Computing, CLOUD* (pp. 983–986). IEEE. <https://doi.org/10.1109/CLOUD.2018.00151>
- Hussien, H., & Aboelnaga, H. (2013). Design of a secured e-voting system. In *2013 International Conference on Computer Applications Technology (ICCAT)* (pp. 1–5). IEEE. <https://doi.org/10.1109/ICCAT.2013.6521985>
- Huszt, A. (2011). A homomorphic encryption-based secure electronic voting scheme. *Publ. Math. Debrecen*, 79(3–4), 479–496. <https://doi.org/10.5486/PMD.2011.5142>
- Jabbar, I., & Alsaad, S. N. (2017). Design and Implementation of secure remote e-voting system using homomorphic encryption. *Int. J. Netw. Secur.*, 19(5), 694–703.
- Jafar, U., Aziz, M. J. A., & Shukur, Z. (2021). Blockchain for electronic voting system—Review and open research challenges. *Sensors*, 21(17), 5874. <https://doi.org/10.3390/s21175874>
- Jafar, U., Aziz, M. J. A., Shukur, Z., & Hussain, H. A. (2022). A cost-efficient and scalable framework for e-voting system based on Ethereum blockchain. In *2022 International Conference on Cyber Resilience (ICCR)* (pp. 1–6). IEEE. <https://doi.org/10.1109/ICCR56254.2022.9996026>

- Kaliyamurthi, K. P., Udayakumar, R., Parameswari, D., & Mugunthan, S. N. (2013). Highly secured online voting system over network. *Indian Journal of Science and Technology*, 6(6), 1–6. <https://doi.org/10.17485/ijst/2013/v6isp6.15>
- Kshetri, N., & Voas, J. (2018). Blockchain-Enabled E-Voting. *IEEE Software*, 35(4), 95–99. <https://doi.org/10.1109/MS.2018.2801546>
- Mursi, M. F. M., Assassa, G. M. R., Abdelhafez, A., & Samra, K. M. A. (2013). On the development of electronic voting: a survey. *International Journal of Computer Applications*, 61(16). <https://doi.org/10.5120/10009-4872>
- Nofer, M., Gomber, P., Hinz, O., & Schiereck, D. (2017). Blockchain. *Business & Information Systems Engineering*, 59(3), 183–187. <https://doi.org/10.1007/s12599-017-0467-3>
- Ogburn, M., Turner, C., & Dahal, P. (2013). Homomorphic Encryption. *Procedia Computer Science*, 20, 502–509. <https://doi.org/10.1016/j.procs.2013.09.310>
- Osgood, R. (2016). The future of democracy: Blockchain voting. *COMP116: Information Security*, 1–21.
- Probor, M. N., Ahmed, M., Kabir, S. B., Fuad, M. M., & Bushra, T. (2023). *Blockchain based e-voting system with homomorphic encryption and threshold signature*. Brac University.
- Qu, W., Wu, L., Wang, W., Liu, Z., & Wang, H. (2022). An electronic voting protocol based on blockchain and homomorphic signcryption. *Concurrency and Computation: Practice and Experience*, 34(16). <https://doi.org/10.1002/cpe.5817>
- Ravindran, S., & Kalpana, P. (2013). *Data Storage Security Using Partially Homomorphic Encryption in a Cloud*. <https://api.semanticscholar.org/CorpusID:62933905>
- Rivest, R. L., Shamir, A., & Adleman, L. (1978). A method for obtaining digital signatures and public-key cryptosystems. *Communications of the ACM*, 21(2), 120–126. <https://doi.org/10.1145/359340.359342>
- Sayeed, S., & Marco-Gisbert, H. (2019). Assessing blockchain consensus and security mechanisms against the 51% attack. *Applied Sciences*, 9(9), 1788. <https://doi.org/10.3390/app9091788>
- Smith, A. D., & Clark, J. S. (2005). Revolutionising the voting process through online strategies. *Online Information Review*, 29(5), 513–530. <https://doi.org/10.1108/14684520510628909>
- Stephen, R., & Alex, A. (2018). A review on blockchain security. *IOP Conference Series: Materials Science and Engineering*, 396(1), 12030.
- Tebaa, M., El Hajji, S., & El Ghazi, A. (2012). Homomorphic encryption method applied to Cloud Computing. In *2012 National Days of Network Security and Systems* (pp. 86–89). IEEE. <https://doi.org/10.1109/JNS2.2012.6249248>
- Thakur, S., Olugbara, O. O., Millham, R., Wesso, H. W., & Sharif, M. (2014). Transforming voting paradigm - The shift from inline through online to mobile voting. In *2014 IEEE 6th International Conference on Adaptive Science & Technology (ICAST)* (pp. 1–7). <https://doi.org/10.1109/ICASTECH.2014.7068115>
- Will, M. A., Nicholson, B., Tiehuis, M., & Ko, R. K. L. (2015). Secure voting in the cloud using homomorphic encryption and mobile agents. In *2015 International Conference on Cloud Computing Research and Innovation (ICCCRI)* (pp. 173–184). IEEE. <https://doi.org/10.1109/ICCCRI.2015.30>
- Wu, D. J. (2012). *Using homomorphic encryption for large scale statistical analysis*. FHE-SI-Report, Univ. Stanford, Tech. Rep. TR-dwu4.
- Xiao, S., Wang, X. A., Wang, W., & Wang, H. (2020). Survey on blockchain-based electronic voting. In H. and M. H. Barolli Leonard and Nishino (Ed.), *Advances in Intelligent Networking and Collaborative Systems* (pp. 559–567). Springer International Publishing. https://doi.org/10.1007/978-3-030-29035-1_54

- Yang, X., Yi, X., Nepal, S., Kelarev, A., & Han, F. (2018). A secure verifiable ranked choice online voting system based on homomorphic encryption. *IEEE Access*, 6, 20506–20519. <https://doi.org/10.1109/ACCESS.2018.2817518>
- Yi, X., & Okamoto, E. (2013). Practical Internet voting system. *Journal of Network and Computer Applications*, 36(1), 378–387. <https://doi.org/10.1016/j.jnca.2012.05.005>
- Zhao, Z. (2014). An efficient anonymous authentication scheme for wireless body area networks using elliptic curve cryptosystem. *Journal of Medical Systems*, 38, 1–7. <https://doi.org/10.1007/s10916-014-0013-5>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Determinants of Halal Food Purchasing Decisions Among Undergraduates at UiTM Tapah

Ilya Zulaikha Zulkifli^{1*}, Nurul Husna Jamian, Samsiah Abdul Razak, Ahmad Nur Azam Ahmad Ridzuan

Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Perak Branch, Tapah Campus, 35400 Tapah Road, Perak, Malaysia.

ARTICLE INFO

Article history:

Received 20 June 2025

Revised 5 August 2025

Accepted 7 August 2025

Published 1 September 2025

Keywords:

Halal Food
Muslim Consumers
Knowledge
Perception
Attitude
Religiosity

DOI:

10.24191/jcrinn.v10i2.543

ABSTRACT

This study explores the factors influencing halal food purchasing decisions among undergraduates at Universiti Teknologi MARA Perak Branch, Tapah Campus. With the rising demand for halal products among Muslim consumers, understanding these drivers is vital for business and policy makers. The study examines the roles of religiosity, knowledge, perception, awareness, attitude, and branding and promotion in shaping student's purchasing decisions. A quantitative research design was employed, utilizing a structured questionnaire distributed to conveniently selected 104 undergraduates from various programs. The study focused on six determinants such as Religiosity, Knowledge, Perception, Awareness, Attitude, and Brand and Promotion Influence. Meanwhile, the dependent variable is Halal Food Purchasing Decision. Multiple Linear Regression analysis was conducted using IBM SPSS version 23. The results revealed that three determinants such as religiosity, awareness, and attitude significantly influence halal food purchasing decisions as p-value less than 0.05. Specifically, religiosity ($B = 0.177$) and attitude ($B = 0.332$) showed a positive effect, while awareness had a negative effect ($B = -0.114$). This study provides valuable insights for food manufacturers, marketers, and university administrators aiming to align strategies with young Muslim consumer's needs. It also highlights the importance of educational initiatives in enhancing halal knowledge and awareness. Future research could expand the scope by conducting comparative studies across different campuses or student demographics.

1. INTRODUCTION

The global halal food industry has experienced significant growth in recent years, driven by increasing awareness among Muslim consumers regarding religious dietary requirements and the rising demand for safe and high quality food (Ghazali et al., 2022). Halal, an Arabic term meaning "permissible" or "lawful", encompasses not only the ingredients used but also the methods of preparation, processing, and handling in accordance with Islamic law. In Malaysia, where Muslims represent the majority population, the halal

^{1*} Corresponding author. *E-mail address:* ilya2177@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.543>

food sector plays a crucial role in the national economy and consumer lifestyle (Faisal et al., 2024). Among the young Muslim demographic, particularly university students, halal food consumption reflects not only of religious obligation but also reflects personal values, awareness, and social influences (Ambali & Bakar, 2014). However, despite the widespread availability of halal-certified products, there remains a gap in understanding the factors influencing undergraduates' purchasing decisions, particularly in campus settings.

At Universiti Teknologi MARA Perak Branch, Tapah Campus, a diverse student population presents a unique opportunity to explore this phenomenon. This diversity includes students from various ethnic backgrounds, socioeconomic statuses, religious commitments, and academic disciplines, all of which contribute to a rich and varied perspective on halal food choices. While previous studies have examined halal consumption patterns on a broader scale, limited research has been conducted to specifically address the determinants influencing the purchasing decisions of students within a campus environment. This gap raises important questions about the role of individual religious values, halal knowledge, awareness, perceptions, attitudes, and the impact of branding and promotional strategies on students' purchasing decisions. The problem lies in the lack of empirical data concerning how these factors interact and which of them are most influential in shaping students' halal food choices. This lack of understanding may hamper businesses and food providers may struggle to meet the expectations and needs of this important consumer segment, potentially affecting both customer satisfaction and sales performance.

Therefore, this study aims to investigate the determinants of halal food purchasing decisions among undergraduates at Universiti Teknologi MARA Perak Branch, Tapah Campus. The specific objectives are: (1) to identify which of the religiosity, knowledge, awareness, perception, attitude, and branding influence halal food purchasing decisions; and (2) to determine which of these factors employs the greatest influence. Through this research, it is hoped that valuable insights will be provided to stakeholders in the food industry, educational institutions, and policymakers to better cater to the halal consumer market among Malaysian youth.

2. LITERATURE REVIEW

The halal food market has gained significant attention from both scholars and practitioners due to the rising demand among Muslim consumers globally (Susanty et al., 2025). Numerous studies have been conducted to understand the factors influencing halal food purchasing behaviour, particularly within the youth and student demographics. This section reviews six sources of determinants for halal food purchasing behaviour, their relevance to the current study, and the gaps that this study aims to address.

Religiosity plays a central role in shaping Muslim's food consumption patterns (Owais et al., 2024). According to (Demirel & Yasarsoy, 2017), religious commitment is one of the strongest predictors of halal food consumption. Their study found that individuals with a high degree of religious observance are more likely to prioritize halal certification and adherence to Islamic dietary laws. The relevance of this finding to the current study is significant, as UiTM Tapah students are predominantly Muslim and their purchasing decisions may be strongly influenced by their religious values. However, the study by Demirel and Yasarsoy (2017) focused on a broader adult Muslim population, overlooking student-specific factors like peer influence or campus style. This research seeks to fill that gap by focusing on undergraduates, whose values may differ due to peer influence and lifestyle factors.

Knowledge and awareness of halal principles are critical in shaping consumer attitudes and behaviours. A study by Ismail (2025) emphasized that consumers with greater knowledge of halal standards leads to more discerning purchasing decisions, with awareness campaigns enhancing trust in halal-certified products. For undergraduates, especially those in non-religious academic fields, the level of halal knowledge may vary considerably. This study builds upon his work by investigating the knowledge levels

among Universiti Teknologi MARA Perak Branch, Tapah Campus students and how it affects their purchasing decisions.

Consumer's attitude and perceptions have been widely studied as predictors of halal purchasing behaviour (Shalihin & Alda, 2025; Youn et al., 2021). Shalihin and Alda (2025) explored the Theory of Planned Behaviour (TPB) in the context of halal food consumption and found that attitude significantly influences purchase intention. Their study demonstrated that a positive attitude toward halal products, combined with perceived behavioural control and subjective norms, increases the likelihood of purchasing halal items. However, their research targeted working adults in urban areas, where purchasing power and social influences differ from those of students. This research extends TPB to a campus context, investigating how attitudes and perceptions shape students' halal food choices.

Branding and promotion strategies also play a role in halal food purchasing decisions. According to Junita et al. (2025), branding influences consumer trust and loyalty, especially when halal certification is prominently displayed on packaging. The study concluded that marketing communication, including social media promotions and halal labels, significantly affects purchasing behaviour among youth. This is particularly relevant for university students, who are heavy users of digital platforms. Nonetheless, the study lacked a specific focus on university campuses, and thus, the current research aims to analyze how branding and marketing appeal to the student population at Universiti Teknologi MARA Perak Branch, Tapah Campus, where product availability and marketing exposure may differ from urban settings due to its semi-rural location. In Universiti Teknologi MARA Perak Branch, Tapah Campus, students may face limited access to a variety of branded halal food products compared to urban campuses, and exposure to aggressive marketing campaigns such as billboards or in-store promotions may also be lower. As a result, students here might rely more on online information or word-of-mouth recommendations when making halal food purchasing decisions.

3. METHODOLOGY

Primary data were utilized to examine the influence of several determinants on the halal food purchasing decision. The target population of this study is all students who are enrolled at Universiti Teknologi MARA (UiTM), Perak Branch, Tapah Campus. The study involved 104 students from Universiti Teknologi MARA Perak Branch, Tapah Campus, who were conveniently selected from various programs. The information was collected using an online questionnaire constructed in Google Forms and distributed via WhatsApp and Telegram platforms. The study focused on six determinants such as religiosity, knowledge, perception, awareness, attitude, and brand and promotion influence. Meanwhile the dependent variable is Halal Food Purchasing Decision. Multiple Linear Regression analysis was conducted using IBM SPSS version 23. Fig. 1 presents the theoretical framework adapted from (Amalia et al., 2020; Chong et al., 2022; Harun et al., 2023) in this study. The assumptions underlying the regression model are discussed in the findings section.



Fig. 1: Theoretical framework of the study

Multiple linear regression analysis is a statistical method used to examine the relationship between one dependent variable and multiple independent variables simultaneously (Zulkifli et al., 2019). Prior to the

analysis, the assumptions of multiple linear regression are assessed to ensure validity. The hypotheses tested in this study are as follows:

H₀: Religiosity, Knowledge, Perception, Awareness, Attitude, and Brand and Promotion Influence have no significant influence on the Halal Food Purchasing Decision.

H₁: Religiosity, Knowledge, Perception, Awareness, Attitude, and Brand and Promotion Influence have a significant influence on the Halal Food Purchasing Decision.

The null hypothesis (H₀) will be rejected if the p-value is less than or equal to the significance level ($\alpha = 0.05$) as a study by Zulkifli et al., (2019), indicating that one or more variables significantly influence the halal food purchasing decision. The multiple linear regression model used in this analysis is presented in Equation (1).

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \dots + \beta_n X_n + \varepsilon \quad (1)$$

where Y represents the dependent variable, $X_1, X_2, \dots, X_n$ denote the determinants, ε is the residual term, and the β 's coefficients represent the regression coefficients with β_0 being the intercept (constant term). There are four assumptions of the regression analysis must be satisfied in order to make the analysis reliable and valid. The assumptions are as follow:

- a) The values of the residuals are normally distributed.
- b) The values of the residuals are independent.
- c) No multicollinearity exists.
- d) There is no outlier exist in dependent variable

4. FINDINGS

A) Assumptions of Multiple Linear Regression

(i) Results of normality test of the residuals

A normal P-P plot of the model residuals was constructed to assess the normality assumption, as shown in Fig 2. The residuals are considered normally distributed when the data points closely follow the diagonal line (Abid & Savikri, 2025). In this case, the assumption of normality is satisfied as the points are closely aligned with the diagonal line.

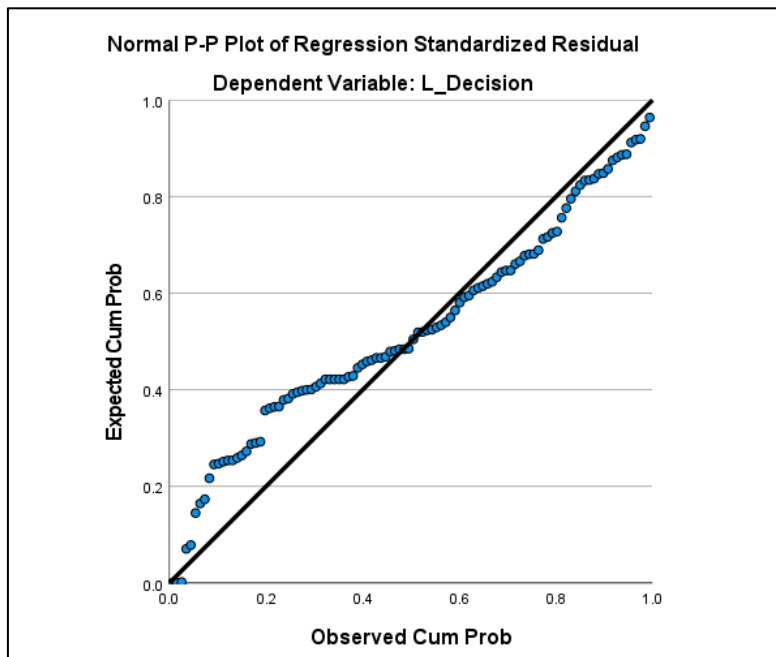


Fig. 2. The normal P-P plot for the residuals of the model

(ii) Results of the independent for the residuals

The Durbin-Watson statistic was used to test the assumption that residuals are independent or uncorrelated. In this study, the Durbin-Watson value was 1.786, which falls within the acceptable range of 1 to 3 referring a study by Zulkifli et al., (2019). This indicates that the assumption of independence among the residuals is met.

(iii) Results of multicollinearity test

The collinearity statistics presented in Table 1 were used to assess the presence of multicollinearity. The analysis indicated that this assumption was satisfied, as all Variance Inflation Factor (VIF) values were below 10 and tolerance values exceeded 0.2 (Zulkifli et al., 2019). Therefore, multicollinearity is not a concern in this dataset.

Table 1. Collinearity Statistics

Variables	Collinearity Statistics	
	Tolerance	VIF
X ₁	0.387	2.584
X ₂	0.326	3.067
X ₃	0.456	2.191
X ₄	0.482	2.076
X ₅	0.409	2.445
X ₆	0.469	2.133

(iv) Results of the existence for the outliers

As shown in Table 2, all Cook's Distance values were below 1 (Zulkifli et al., 2019), indicating that no individual cases exert an undue influence on the model. This suggests the absence of significant outliers that could affect the regression results.

Table 2. Cook's Distance Value

Cook's Distance Statistics	Statistics
Minimum	0.6213
Maximum	0.7145
Mean	0.6891
Standard Deviation	0.0173
Size (N)	104

(v) Multiple linear regression analysis

The F-statistic obtained was 13.088 with a p -value less than 0.05 indicating that the overall regression model is statistically significant at the five percent significance level. As presented in Table 3, the regression results reveal that Religiosity, Awareness, and Attitude are significant determinants at the 0.05 level. This suggests that these three determinants have a meaningful influence on students' decisions to purchase halal food. The estimated coefficients for Religiosity, Awareness, and Attitude are 2.876, -3.807, and 5.104, respectively. Therefore, the null hypothesis is rejected, confirming that Religiosity, Awareness, and Attitude significantly influence halal food purchasing decisions among students. The regression equation is presented in Equation (2).

$$\hat{Y} = 8.984 + 2.876X_1 - 3.807X_4 + 5.104X_5 \quad (2)$$

The 2.876 associated with halal food purchasing decision indicates that for each additional in level of religiosity, the decision level will increase 2.876, holding the other determinants constant. Then for each additional in Awareness, the decision level will decrease 3.807, if the other determinants remain constant. Lastly for Attitude, every one level increase in attitude, the halal food purchasing decision increases by 5.104 with all other determinants held constant.

Table 3. Regression analysis results

Model	B	Standard Error	t	p-value
β_0	0.380	0.042	8.984	0.000*
X_1	0.177	0.062	2.876	0.005*
X_2	0.059	0.079	0.745	0.458
X_3	0.027	0.072	0.370	0.712
X_4	-0.114	0.030	-3.807	0.000*
X_5	0.332	0.065	5.104	0.000*
X_6	-0.019	0.037	-0.525	0.601

Note: * p -value ≤ 0.05

5. CONCLUSION AND RECOMMENDATION

This study investigated the influence of six determinants Religiosity, Knowledge, Perception, Awareness, Attitude, and Brand and Promotion Influence on students' decisions to purchase halal food at Universiti Teknologi MARA Perak Branch, Tapah Campus. The analysis employed multiple linear regression to examine these relationships, with all regression assumptions being satisfactorily met, including normality, independence of residuals, absence of multicollinearity, and lack of influential outliers. The results revealed that three determinants such as religiosity, awareness, and attitude significantly influence halal food purchasing decisions. Specifically, religiosity and attitude showed a positive effect, while awareness had a negative effect. The statistical significance of the model, supported by an F-statistic of 13.088 and a *p*-value below 0.05, confirms the overall validity of the regression model.

The insights provided may guide marketers, educators, and policymakers in developing more targeted strategies to promote halal food consumption especially within educational institutions. Future research could explore additional variables or apply the model to broader demographic groups, use interviews or focus groups, and explore the role of social media, peer influence, and lifestyle trends. Long-term studies could also track changes over time. These steps can give a deeper and broader understanding of halal food choices among students. This study is limited by its small sample size and focus on one campus, which may not reflect students at other universities. It also used a quantitative approach, which doesn't capture personal experiences.

6. ACKNOWLEDGEMENT/FUNDING

The author would like to acknowledge the support of Universiti Teknologi MARA Perak Branch, Tapah Campus and Faculty of Computer and Mathematical Sciences for the providing the facilities on this research.

7. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

8. AUTHORS' CONTRIBUTIONS

Nurul Husna Jamian: Conceptualisation, methodology, formal analysis, investigation and writing-original draft; **Ilya Zulaikha Zulkifli:** Conceptualisation, methodology, and formal analysis; **Ahmad Nur Azam Ahmad Ridzuan:** Conceptualisation, formal analysis, and validation; **Samsiah Abdul Razak** Conceptualisation, writing- review and editing and formatting.

9. REFERENCES

- Abid, M., & Savikri, H. (2025). The influence of motivation and work discipline on employee performance. *Golden Ratio of Data in Summary*, 5(2), 320–332. <https://doi.org/10.52970/grdis.v5i2.1200>
- Amalia, F. A., Sosianika, A., & Suhartanto, D. (2020). Indonesian millennials' halal food purchasing: Merely a habit? *British Food Journal*, 122(4), 1185–1198. <https://doi.org/10.1108/BFJ-10-2019-0748>
- Ambali, A. R., & Bakar, A. N. (2014). People's awareness on halal foods and products: Potential issues for policy-makers. *Procedia - Social and Behavioral Sciences*, 121(September 2012), 3–25. <https://doi.org/10.1016/j.sbspro.2014.01.1104>

<https://doi.org/10.24191/jcrinn.v10i2.543>

- Chong, S. C., Yeow, C. C., Low, C. W., Mah, P. Y., & Tung, D. T. (2022). Non-Muslim Malaysians' purchase intention towards halal products. *Journal of Islamic Marketing*, 13(8), 1751–1762. <https://doi.org/10.1108/JIMA-10-2020-0326>
- Demirel, Y., & Yasarsoy, E. (2017). Exploring consumer attitudes towards halal. *Journal of Tourismology*, 3(1), 34–43. <https://doi.org/10.26650/jot.2017.3.1.0003>
- Faisal, A., Adzharuddin, N. A., & Raja Yusof, R. N. (2024). The nexus between advertising and trust: Conceptual review in the context of halal food, Malaysia. *International Journal of Academic Research in Business and Social Sciences*, 14(1), 600–620. <https://doi.org/10.6007/ijarbss/v14-i1/19718>
- Ghazali, E. M., Mutum, D. S., Waqas, M., Nguyen, B., & Ahmad-Tarmizi, N. A. (2022). Restaurant choice and religious obligation in the absence of halal logo: A serial mediation model. *International Journal of Hospitality Management*, 101(January 2021), 103109. <https://doi.org/10.1016/j.ijhm.2021.103109>
- Harun, N. H. binti, Idris, N. A. Z., & Mohamed Bashir, A. (2023). Factors influencing halal food products purchasing among young adults according to theory of planned behavior. *International Journal of Academic Research in Business and Social Sciences*, 13(1). <https://doi.org/10.6007/ijarbss/v13-i1/14809>
- Ismail, I. J. (2025). Halal brand quality and halal food purchasing intention among university students: The moderating effect of customer-employee interactions. *Social Sciences and Humanities Open*, 11(February 2024), 101352. <https://doi.org/10.1016/j.ssaho.2025.101352>
- Junita, S., Nuraeni, I., Munasib, M., Proverawati, A., Ramadhan, R., & Soedirman, U. J. (2025). The relationship between perception of halal certification and purchase decision of contemporary drinks among students of Jenderal Soedirman. *Matan: Journal of Islam and Muslim Society*, 7(1), 46–56. <https://doi.org/10.20884/1.matan.2025.7.1.14696>
- Owais, U., Patel, R., & Abbott, S. (2024). The influence of religiosity on food choice among British Muslims: A qualitative study. *Nutrition and Health*, 1–8. <https://doi.org/10.1177/02601060241244883>
- Shalihin, A., & Alda, T. (2025). Consumer intentions to purchase eco-friendly halal food in Medan, Indonesia: An approach using the theory of planned behavior[†]. In *the Proceedings of The 8th Mechanical Engineering, Science and Technology International Conference* (pp. 1-12). MDPI. <https://doi.org/10.3390/engproc2025084083>
- Susanty, A., Puspitasari, N. B., Jie, F., Akhsan, F. A., & Jati, S. (2025). Consumer acceptance of halal food traceability systems: A novel integrated approach using modified UTAUT and DeLone & McLean models to promote sustainable food supply chain practices. *Cleaner Logistics and Supply Chain*, 15(October 2023), 100226. <https://doi.org/10.1016/j.clscn.2025.100226>
- Youn, H., Xu, J. (Bill), & Kim, J. H. (2021). Consumers' perceptions, attitudes and behavioral intentions regarding the symbolic consumption of auspiciously named foods. *International Journal of Hospitality Management*, 98(December 2020), 103024. <https://doi.org/10.1016/j.ijhm.2021.103024>
- Zulkifli, I. Z., Jamian, N. H., & Zulkipli, F. (2019). Factors influencing environmental sanitation among undergraduates at UiTM, Tapah Campus. *International Journal of Recent Technology and Engineering*, 8(2 Special Issue 11), 598–601. <https://doi.org/10.35940/ijrte.B1092.0982S1119>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Cyberpark: An-IoT based Automated Parking Management and Monitoring System using RFID and ESP32

Nur Amalina Muhamad^{1*}, Norhalida Othman², Nor Diyana Md Sin³, Noor Haliza Abdul Rahman⁴, Muhammad Iskandar Mohd Rodzi⁵

^{1,2,3,4,5}Faculty of Electrical Engineering, Universiti Teknologi MARA Johor Branch, Pasir Gudang Campus, 81750 Masai, Johor, Malaysia.

ARTICLE INFO

Article history:

Received 24 June 2025

Revised 5 August 2025

Accepted 6 August 2025

Online first

Published 1 September 2025

Keywords:

Smart Parking System

Internet of Things (IoT)

RFID

Real-Time Monitoring

ESP32 and Arduino Integration

DOI:

10.24191/jcrinn.v10i2.547

ABSTRACT

The growing number of vehicles in urban and institutional areas has created significant challenges in parking space management. Traditional parking systems, which rely on manual processes, are often inefficient, time-consuming, and lack real-time information, leading to increased traffic congestion, fuel consumption, and user dissatisfaction. To address these issues, this paper presents Cyberpark, an IoT-based Automated Parking Management and Monitoring System that integrates smart technologies to improve parking operations, enhance user convenience, and increase infrastructure safety. Cyberpark utilizes RFID technology for secure and automated vehicle access, infrared (IR) sensors to detect parking slot occupancy, servo motors to control entry and exit gates, and a water sensor to detect flooding within the parking area. The system incorporates two microcontrollers: Arduino UNO for local control and ESP32 for wireless communication with the Adafruit IO cloud platform. Real-time data from all sensors is uploaded to the cloud, allowing for remote monitoring and status updates via an online dashboard. The development process involved circuit simulation using Proteus 8 Professional and functional testing through the Wokwi simulator. Hardware components were assembled on a custom PCB, and the software was developed using the Arduino IDE. Results showed effective system performance in automating gate operations, accurately monitoring parking slot status, and issuing timely alerts during flood conditions. The real-time dashboard allowed for seamless user interaction and remote oversight. Cyberpark demonstrates the feasibility of integrating low-cost, open-source hardware with cloud-based IoT platforms to create a scalable, user-friendly, and environmentally responsive smart parking solution. This system is suitable for deployment in campuses, shopping complexes, or residential areas. The project contributes to the development of smart city infrastructure by offering a reliable, efficient, and safe approach to modern parking management.

1. INTRODUCTION

The increasing number of vehicles in urban areas has escalated the demand for efficient and intelligent parking solutions. Traditional parking systems, which depend on manual management or basic mechanical

^{1*} Corresponding author. E-mail address: amalina0942@uitm.edu.my

infrastructure, often struggle with inefficiencies such as poor space utilization, long vehicle queues, and driver frustration. These inefficiencies result in increased fuel consumption and carbon emissions as drivers circulate in search of available parking spaces, contributing to environmental degradation and worsening urban traffic congestion (El Bilali et al., 2022; Suhartono et al., 2023).

Smart parking systems have emerged as a transformative solution, leveraging the capabilities of the Internet of Things (IoT), embedded sensors, microcontrollers, and wireless communication to address modern urban challenges. These systems aim to provide real-time monitoring of parking slot occupancy, automate gate control, enhance user experience, and improve operational efficiency. When integrated with cloud platforms, smart parking systems can also enable remote monitoring, data analytics, and predictive maintenance, aligning with broader smart city initiatives (Alymani et al., 2025).

The deployment of microcontrollers such as ESP32 and Arduino UNO in smart parking solutions has gained popularity due to their cost-effectiveness, programmability, and Wi-Fi connectivity. Coupled with RFID (Radio Frequency Identification) for access control, infrared (IR) or ultrasonic sensors for presence detection, and cloud services like Adafruit IO or Firebase for data visualization, these systems offer real-time visibility and automation (Ahmed et al., 2018; Gowda et al., 2021; Vantagudi et al., 2024). Furthermore, incorporating environmental sensors—like flood detectors—addresses safety concerns in flood-prone regions, offering added value to smart parking infrastructures (Abdulsahab et al., 2020; Alsafar et al., 2025).

Numerous studies have demonstrated the potential of emerging technologies to enhance the efficiency and functionality of parking management systems. For example, Azmi and Ismail (2021) utilized ultrasonic sensors in conjunction with a Raspberry Pi to improve the speed and accuracy of parking slot detection, thereby reducing vehicle idle time. In a similar vein, Mani and Bhat (2020) developed an ESP8266-based system incorporating infrared sensors and LCD interfaces, which contributed to alleviating traffic congestion within parking facilities. Expanding on these foundational models, Gowda et al. (2021) implemented RFID-enabled slot allocation using the ESP32 microcontroller, streamlining the entry and exit processes through automated access control. Additionally, Elakya et al. (2019) proposed a hybrid framework that integrated ESP32, RFID, and IR sensors, thereby enhancing the responsiveness and adaptability of parking operations. To address environmental concerns, Abdulsahab et al. (2020) incorporated flood detection capabilities into their Adafruit IO-based system, offering real-time alerts in adverse weather conditions and strengthening the system's safety features. Building upon these innovations, Alymani et al. (2025) introduced machine learning algorithms for occupancy prediction in IoT-based smart parking environments, thereby increasing the system's predictive accuracy and operational efficiency. Similarly, Alsafar et al. (2025) designed an advanced IoT-integrated solution focused on mitigating traffic congestion and reducing vehicular emissions, aligning with broader sustainability goals. Despite these advancements, existing smart parking systems continue to exhibit limitations, including the lack of comprehensive environmental integration, challenges in scalability, insufficient user feedback mechanisms, and reliance on centralized infrastructure. These gaps underscore the necessity for a more cohesive, adaptable, and user-oriented smart parking architecture.

Therefore, this paper introduces Cyberpark: Automated Parking Management and Monitoring System, a cost-effective, modular, and user-friendly solution that seamlessly integrates RFID-based access control, servo motor-driven gate operation, real-time IR occupancy detection, flood monitoring, and cloud-based IoT connectivity into a unified architecture. By synthesizing these features into a single scalable platform, Cyberpark contributes to the field of smart infrastructure, enhancing urban mobility, improving safety, and promoting sustainability.

2. METHODOLOGY

The development process for the **Cyberpark: Automated Parking Management and Monitoring System** was carried out in three main phases: simulation, hardware construction, and cloud integration. The

<https://doi.org/10.24191/jcrinn.v10i2.547>

initial stage emphasized simulation and design validation using tools such as **Proteus 8 Professional** and **Wokwi**. These simulations enabled the design and testing of circuit schematics, system logic, and component interactions within a virtual environment. Concurrently, program codes for both **Arduino UNO** and **ESP32** microcontrollers were developed and tested using the **Arduino IDE**, ensuring functional integrity before physical implementation. After completing the simulation phase, the project advanced to hardware prototyping. This phase involved assembling electronic components, wiring connections on a breadboard, and finally transferring the validated layout onto a custom-made **printed circuit board (PCB)**. The core of the system includes **Arduino UNO** for hardware control (RFID, servo motors, water sensor, and buzzer) and **ESP32** for IoT communication. Key modules used are **MFRC522** RFID readers, **infrared sensors** for parking slot detection, **servo motors** for gate automation, **water level sensor**, **LED indicators**, and an **LCD display** for user feedback. The choice of Arduino UNO was due to its suitability for handling input/output control, while ESP32 was selected for its built-in Wi-Fi functionality and compatibility with real-time cloud systems. In the final stage, the ESP32 was configured to interface with the **Adafruit IO** cloud platform using the **MQTT protocol**. This connection allows real-time data transmission for remote monitoring of parking occupancy and flood alerts. The system sends continuous updates to the cloud, enabling users and operators to access live data via an intuitive dashboard, either on a web browser or mobile device.

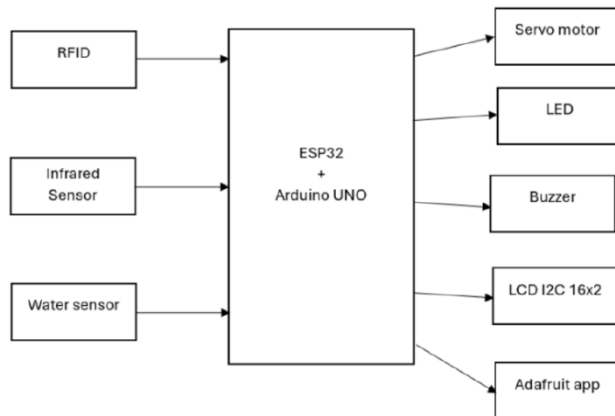


Fig. 1. Block diagram of Cyberpark: Automated parking management and monitoring system

Fig. 1. illustrates the block diagram of the Cyberpark: Automated Parking Management and Monitoring System, showcasing the interaction between input devices, processing units, and output components. The system's core is built on the integration of ESP32 and Arduino UNO, which serve as the main processing units to control operations and handle communication. The input section consists of three key sensors: the RFID module, infrared (IR) sensor, and water level sensor. The RFID module authenticates vehicles at the entry and exit points. The IR sensors detect vehicle presence in parking slots, allowing the system to determine whether a slot is occupied or vacant. The water sensor monitors flooding conditions within the parking area. Based on these inputs, the system activates various output components. The servo motors open or close the gates when valid RFID cards are scanned. LED indicators display the status of each parking space (red for occupied, green for available). The buzzer alerts users during flood detection, while the LCD 16x2 displays real-time instructions and status. Additionally, sensor data is sent via ESP32 to the Adafruit cloud dashboard, enabling remote monitoring and data logging. This integration enables a responsive, safe, and efficient smart parking environment.

Fig. 2. illustrates the operational flow of the Cyberpark: Automated Parking Management and Monitoring System, outlining the sequence of actions involved in managing vehicle entry, parking

occupancy, environmental monitoring, and vehicle exit. The process begins when a car arrives at the entrance. The driver scans an RFID card, which is verified by the system. If the card is valid, the servo motor opens the gate, allowing access. If invalid, entry is denied. Upon successful entry, the parking slot is marked as occupied, and the red LED is activated to indicate the slot is taken. Simultaneously, the system continuously monitors the environment using a water level sensor. If the water level exceeds the threshold (e.g., level > 10), a buzzer is triggered to alert users to potential flooding. If no risk is detected, the system proceeds with regular monitoring. As the car parks, infrared sensors detect its presence and update the parking status. When a space is occupied, the red LED turns on; if vacant, the green LED is illuminated. The status of each slot is sent in real-time to the Adafruit IO cloud platform via ESP32, enabling remote monitoring through a dashboard. When the driver is ready to exit, the RFID card is scanned again at the exit gate. Upon successful verification, the servo motor opens the gate, and the slot is marked as available. The system updates the LED status and sends new data to the cloud. This flowchart summarizes the integration of hardware and software to deliver a smart, automated parking solution that improves efficiency, enhances safety, and provides real-time data to users and operators.

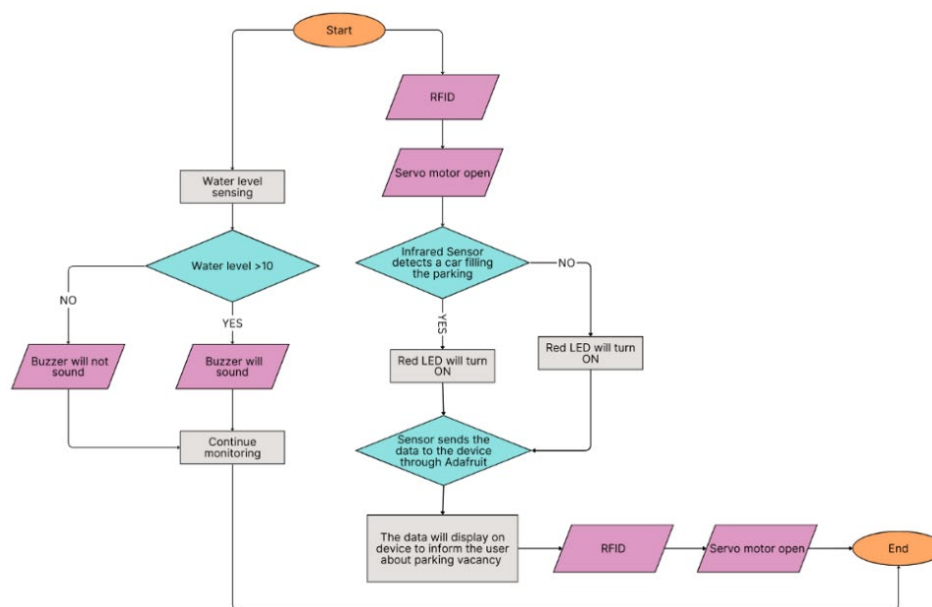


Fig. 2. Flow chart of the Cyberpark

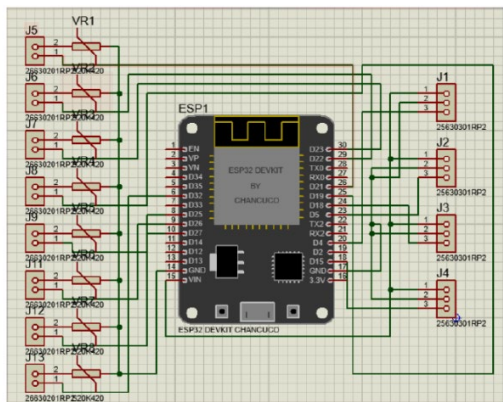


Fig. 3. The schematic diagram of the Cyberpark

<https://doi.org/10.24191/jcrinn.v10i2.547>

Fig. 3. displays the schematic wiring diagram of the Cyberpark system, built around the ESP32 DevKit microcontroller. This diagram illustrates the connections between the ESP32 and various input/output components through multiple digital and analog GPIO pins. The diagram shows how different devices such as infrared sensors, RFID readers, water sensors, LEDs, buzzers, and LCD displays are interfaced with the microcontroller for system functionality. On the left side, a series of components such as IR sensors and input devices are connected to digital pins D2 through D15. These pins handle signals for parking space detection, RFID card scanning, and water level monitoring. On the right side, the diagram shows output devices connected to pins D16 through D30, which include the LCD, servo motor, LED indicators, buzzer, and cloud communication interface.

This schematic serves as the electronic foundation of the Cyberpark system, ensuring all peripherals are appropriately powered and assigned to correct pins for accurate operation. Table 1 presents the detailed connection mapping between the ESP32 microcontroller and various components used in the Cyberpark system, including sensors, actuators, display units, and communication interfaces.

Table 1. Connection of component with ESP32 microcontroller

Component	Pin Description	ESP32 Pin Connection	Remarks
RFID Module (MFRC522)	SDA	D21	SPI/I2C Communication
	SCK	D18	Clock signal
	MOSI	D23	Master Out Slave In
	MISO	D19	Master In Slave Out
	RST	D22	Reset
Infrared Sensor 1	Signal	D2	Parking Slot Detection
Infrared Sensor 2	Signal	D4	Parking Slot Detection
Water Sensor	Analog Out	VP (A0)	Flood Detection
Servo Motor	PWM Control	D5	Gate Control
Red LED	Positive	D12	Slot Occupied Indicator
Green LED	Positive	D13	Slot Vacant Indicator
Buzzer	Positive	D14	Flood Alert Sound
LCD 16x2 (I2C)	SDA	D21	Shared with RFID (I2C Bus)
	SCL	D22	Shared with RFID (I2C Bus)
Adafruit IO (via WiFi)	-	Built-in WiFi (ESP32)	Cloud Communication



Fig.4. The completed prototype of the Cyberpark

Fig. 4 displays the completed prototype of the Cyberpark: Automated Parking Management and Monitoring System, which visually represents a scaled-down parking facility integrated with smart automation features. The setup includes miniature lanes, servo-controlled entry and exit gates, and parking

<https://doi.org/10.24191/jcrinn.v10i2.547>

bays equipped with infrared (IR) sensors and LED indicators. Two RFID modules are mounted at the gate areas to authenticate vehicles using RFID cards. Green and red LEDs positioned above each parking bay indicate available and occupied slots, respectively, based on real-time sensor input. A buzzer is included to alert users during flood simulation, while an LCD display at the front panel guides users with system messages such as “WELCOME” or “SCAN CARD.” Internally, the ESP32 microcontroller manages cloud communication with the Adafruit IO platform, while the Arduino UNO handles gate control and hardware interfacing. The prototype demonstrates smooth coordination between access control, slot monitoring, and remote data updating, effectively simulating the real-world operation of a smart parking system.

3. RESULTS AND DISCUSSION

3.1 Simulation flow of the Cyberpark system

The simulation of the Cyberpark system was developed using the Wokwi platform to validate the functional flow of hardware components before physical implementation. Each stage of the simulation represents a core part of the parking management operation, including entrance access, vehicle detection, gate automation, and cloud monitoring via IoT. Fig. 5. illustrates the initial state of the entrance system, where users are prompted to input a code through a 4x4 keypad to simulate RFID card authentication. Upon activation, the LCD screen displays “Enter Password!” to prompt the user. The Arduino microcontroller processes the input and prepares to respond based on authentication results.

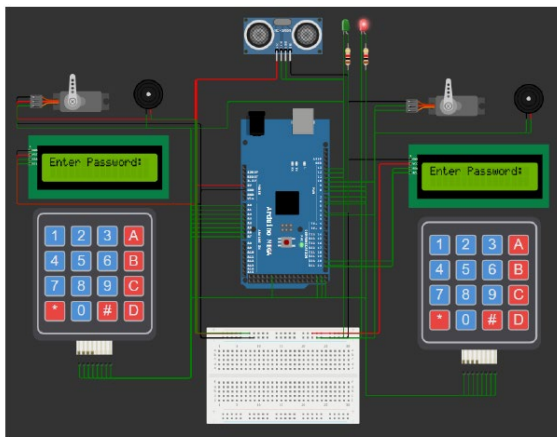


Fig. 5. Simulation diagram of the Cyberpark

Once the correct code is entered, as shown in Fig. 6. (left) the servo motor simulates gate operation by rotating to allow the vehicle to enter. Simultaneously, the LCD screen updates to display “Access Granted,” confirming successful authorization. This mechanism effectively mimics the RFID-based entry control in a real parking system and ensures secure access control. Fig. 6. (right) focuses on the vehicle detection system using the HC-SR04 ultrasonic sensor, which simulates the presence of a vehicle within a parking slot. When a vehicle is detected within a specified range, the red LED turns on, indicating that the slot is occupied. If no vehicle is detected, the green LED lights up, showing the space is vacant. This sensor logic ensures accurate, automated status monitoring of each parking bay.

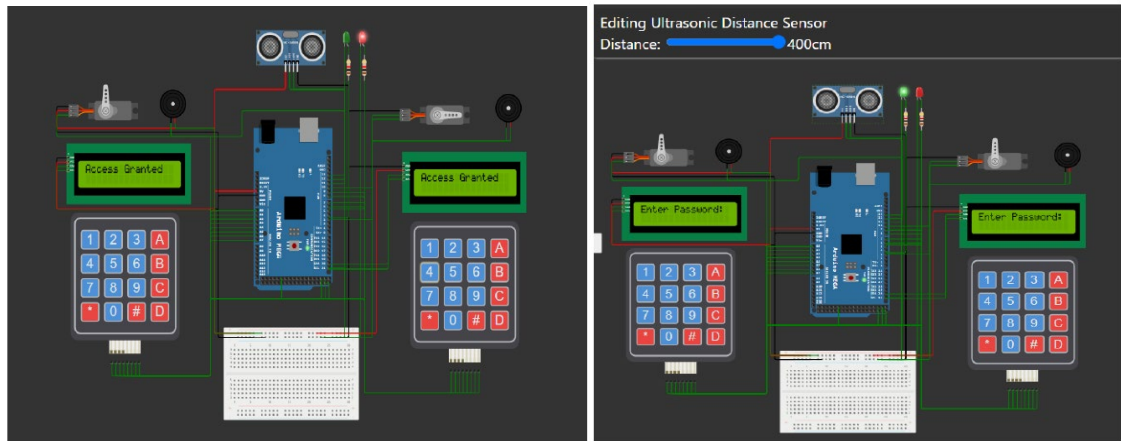


Fig. 6. (left) Simulation of entrance system and Fig. 6. (right) Simulation of sensor detection

The simulation for the IR sensor and LED functionality is represented using the HC-SR04 ultrasonic sensor in Wokwi. This setup detects the presence of a vehicle in the parking space. When a car is detected, the red LED activates, signalling the space is occupied. If no vehicle is detected, the green LED lights up.

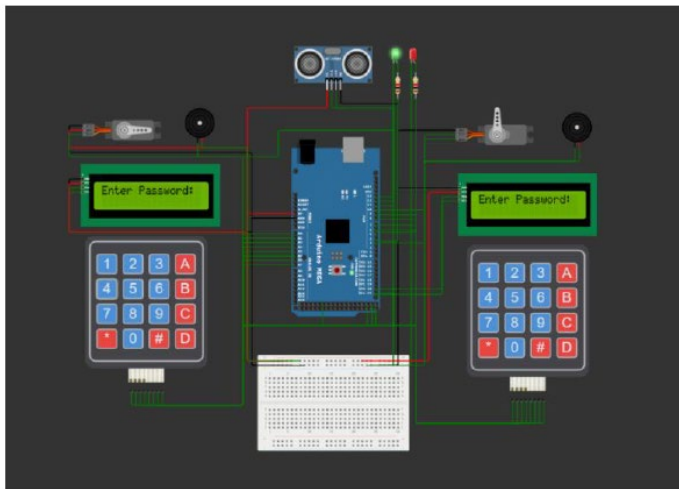


Fig. 7. Simulation of exit process

The exit process is demonstrated in Fig. 7, where a similar keypad is used to simulate RFID verification for vehicle departure. Upon entering the correct exit code, the servo motor reactivates to open the gate, allowing the vehicle to leave. A delay mechanism ensures the gate closes automatically after a short period, simulating safe and controlled exit behaviour.

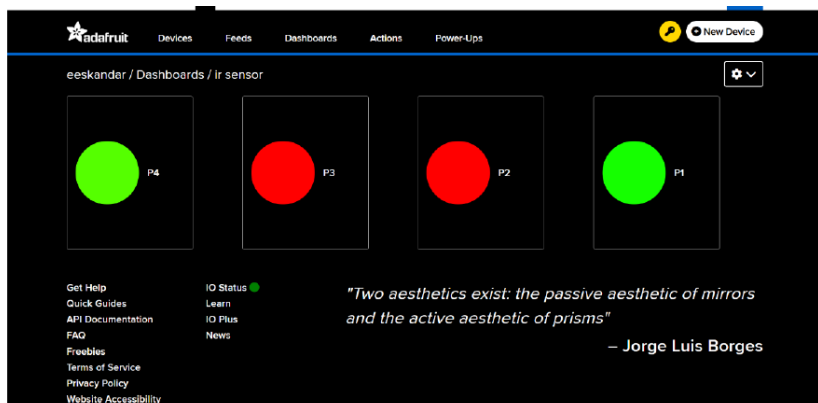


Fig. 8. IoT dashboard of Cyberpark

Finally, Fig. 8. presents the IoT dashboard interface via Adafruit IO, which displays real-time slot status. The dashboard shows circular indicators representing each parking bay, turning green for available and red for occupied. This cloud-based monitoring system allows users or administrators to remotely observe the availability of parking slots, ensuring timely decision-making and improved user experience. The dashboard connectivity highlights the effectiveness of IoT integration in modern smart parking systems. Together, this simulation sequence validates the full operational cycle of the Cyberpark system from secure entry, parking detection, and exit, to live cloud monitoring demonstrating the functionality and reliability of each integrated component.

3.2 Final hardware implementation and testing

Fig. 9. illustrates the vehicle entrance operation in the Cyberpark prototype. The system uses RFID technology for secure access control. As a vehicle approaches the gate, the driver presents an RFID card. The Arduino Uno reads and verifies the card's authenticity. Upon successful validation, a servo motor is triggered to lift the gate, allowing the vehicle to enter. To guide the user, an LCD screen mounted on the prototype displays status messages such as "WELCOME" or "ACCESS GRANTED," enhancing user interaction and ensuring a smooth entry process. This component reflects the real-world scenario of automated entry gates in smart parking facilities.



Fig. 9. Vehicle entrance

Fig. 10. demonstrates how the system uses infrared (IR) sensors to detect the presence of a vehicle once parked inside a bay. Each sensor is connected to the ESP32 microcontroller, which processes occupancy data in real-time. If a vehicle is detected, the corresponding red LED turns on, indicating that the slot is taken. This information is then pushed to the Adafruit IO cloud platform, enabling remote monitoring of parking space availability. When a vehicle exits, the sensor detects the absence and switches the LED to green, signalling that the space is available again. This setup ensures accurate and efficient management of parking slots with minimal human intervention.

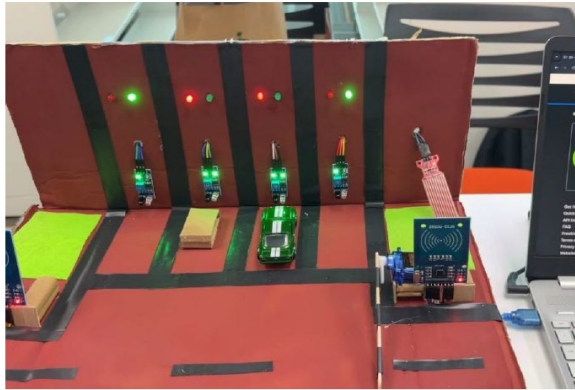


Fig. 10. Sensor detects the vehicle

Fig. 11. shows the water detection system in action. A water level sensor is submerged in a container to simulate flooding. When the sensor detects that the water level has exceeded a critical threshold (analog reading > 500), a buzzer is activated to alert users and system operators of the flood condition. This feature adds an important safety layer, ensuring that environmental hazards such as flooding are actively monitored. This helps prevent vehicle damage during heavy rain and supports the overall resilience of the parking infrastructure. Together, these hardware modules validate the complete flow of the Cyberpark system from access authorization and parking status monitoring to environmental safety demonstrating its readiness for real-world smart parking deployment.



Fig. 11. The water detection as a flood system alert

4. CONCLUSION

This project successfully developed and demonstrated Cyberpark, an IoT-based automated parking management and monitoring system that addresses common challenges in urban and institutional parking facilities. By integrating RFID-based access control, infrared sensors for real-time occupancy detection, servo motor-driven gates, environmental monitoring via water sensors, and cloud connectivity using Adafruit IO, the system offers a comprehensive solution for smart parking operations. The project began with simulations using Proteus and Wokwi, validating the logic and component interactions before hardware implementation. The prototype was built using Arduino UNO and ESP32, combining local hardware control with wireless cloud communication. Results from the hardware testing showed that the system effectively automated vehicle entry and exit, accurately monitored parking slot availability, and reliably detected flood conditions. Real-time feedback was provided through LEDs, LCDs, buzzers, and a cloud dashboard, ensuring both user interaction and remote monitoring capabilities. The prototype proved to be cost-effective, scalable, and user-friendly, making it suitable for deployment in environments such as campuses, residential areas, or commercial buildings. In addition, the inclusion of environmental safety features such as flood alerts adds value beyond traditional parking systems. In conclusion, Cyberpark demonstrates the potential of combining embedded systems and IoT technologies to build intelligent, responsive infrastructure. Future work may include expanding the number of parking bays, integrating license plate recognition, or enabling mobile app access to enhance the overall system performance and user experience.

5. ACKNOWLEDGEMENTS

The authors would like to acknowledge the support of Universiti Teknologi MARA (UiTM, Johor Branch, Pasir Gudang Campus and Faculty of Electrical Engineering for supporting this study.

6. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

7. AUTHOR'S CONTRIBUTION

Nur Amalina Muhamad: Project coordination, manuscript writing, and overall supervision; **Norhalida Othman:** Data analysis, validation of system functionality, and manuscript editing; **Nor Diyana Md Sin:** Literature review, system testing, and result interpretation; **Noor Hasliza Abdul Rahman:** Technical documentation, graphical representation, and formatting; **Muhammad Iskandar Mohd Rodzi:** System design, hardware integration, and prototype development.

8. REFERENCES

- Abdulsahab, J. A., Hussien, A. A., & Al-Naji, L. A. (2020). IoT based smart parking system. *International Journal of Scientific & Engineering Research*, 11(3), 50–55.
- Ahmed, M. M., Chowdhury, S. A., & Reza, M. M. (2018). IoT-based smart parking system for urban cities. *International Journal of Computer Applications*, 180(20), 7–10.

- Alymani, M., Alsharif, M. H., & Alfarraj, O. (2025). Enabling smart parking for smart cities using Internet of Things (IoT) and machine learning. *PeerJ Computer Science*, 11, e2544. <https://doi.org/10.7717/peerj-cs.2544>
- Alsafar, A., Bader, A., & Hamad, S. (2025). Advanced IoT-integrated parking systems with automated license plate recognition and flood detection. *Scientific Reports*, 15(1), 1123. <https://doi.org/10.1038/s41598-025-1123-x>
- Azmi, N. M. F. A. B., & Ismail, M. B. (2021). Smart parking system using IoT with ultrasonic sensor. *National Seminar on Embedded Systems*, 1–5.
- El Bilali, H., Arrar, A., & Cherti, L. (2022). Intelligent parking systems and their economic and environmental impact: A review. *Energy Reports*, 8, 240–251. <https://doi.org/10.1016/j.egy.2022.09.012>
- Elakya, R., Lakshmi, P., & Sharmila, D. (2019). Smart parking system using IoT. *International Research Journal of Engineering and Technology (IRJET)*, 6(4), 275–279.
- Gowda, Y., Ashwini, R., & Kiran, M. (2021). Smart parking system using IoT. *International Journal of Advanced Research in Computer & Communication Engineering*, 10(6), 200–204.
- Mani, R. M., & Bhat, V. V. (2020). Smart parking system. *International Journal of Engineering Research & Technology (IJERT)*, 9(3), 122–125. <https://doi.org/10.17577/IJERTV9IS090305>
- Suhartono, P., Nugraha, H. R., & Pratama, M. D. (2023). Development of smart parking systems using RFID sensors for online verification in different areas. *IoT and Applications (IOTA)*, 3(2), 45–50.
- Vantagudi, X., Sharma, R., & Raj, K. (2024). IoT-enabled smart car parking system through integrated sensors and mobile applications. *arXiv preprint*. <https://arxiv.org/abs/2403.10123>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

A Web-Based System for Managing Student Attendance and Assessment Submissions: An SDLC Waterfall Model Approach

Siti Balqis Mahlan¹, Jamal Othman^{2*}, Maisurah Shamsuddin³, Syarifah Adilah Mohamed Yusoff⁴, Zalina Othman⁵

^{1,2,3,4}Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Pulau Pinang Branch, Permatang Pauh Campus, 13500 Pulau Pinang, Malaysia.

⁵Institut Teknologi Petroleum PETRONAS (INSTEP), Batu Rakit, 21020 Kuala Nerus, Terengganu, Malaysia.

ARTICLE INFO

Article history:

Received 12 June 2025

Revised 30 July 2025

Accepted 31 July 2025

Published 1 September 2025

Keywords:

Student's Attendance

Assignment Submission

SDLC

Web-Based

Cronbach's Alpha

DOI:

10.24191/jcrinn.v10i2

ABSTRACT

University lecturers handle numerous tasks, including tracking student attendance, managing multiple assessment submissions, handling online tests and quizzes, as well as providing timely evaluations of assessments. Currently, these processes are manually managed, which can be time-consuming and inefficient due to the absence of an integrated web-based application system. For example, students often submit their assignments or project papers through cloud storage links provided by lecturers. However, the lack of standardization in file naming frequently forces lecturers to spend additional time renaming hundreds of files before they can begin the grading process. This paper presents a proposed solution in the form of an integrated web-based application system designed to automate and streamline critical tasks like attendance tracking, assignment submission, and evaluation, serving as a comprehensive, one-stop solution for lecturers. SDLC (System Development Life Cycle) Waterfall Development Approach was applied as a method of system construction. HTML and PHP were employed for its implementation and handling of diverse functionalities. To assess the system's effectiveness, it was tested over one academic semester with the participation of 111 degree and diploma students at Universiti Teknologi MARA, Perlis branch. Findings revealed that the system significantly simplified the attendance logging process and streamlined assignment submissions, as perceived by the students. In parallel, lecturers reported that the system facilitated a more efficient grading process, allowing them to return assessments promptly and allowing students to access their marks quickly and easily. Overall, the integration of these automated modules helped reduce bureaucratic inefficiencies, allowing lecturers to perform their duties with greater efficiency and less stress. Positive feedback from students also highlighted a desire for this system to be expanded to other courses, alongside suggestions for further enhancements to improve functionality.

^{2*} Corresponding author. E-mail address: jamalothman@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.539>

1. INTRODUCTION

Lecturers have multifaceted responsibilities that extend beyond mere teaching planning and preparation. Academicians are responsible for accomplishing various duties; for instance, recording student attendance and absenteeism, evaluating assignments, providing consultations, acting as academic advisors, and handling numerous other essential tasks (AGCAS, 2021).

At Universiti Teknologi MARA, lecturers must meticulously record and monitor student attendance over a 14-week class period. To illustrate the workload, imagine a lecturer handling an average of three groups per subject, with two subjects typically assigned. Managing attendance for this many students involves an overwhelming amount of administrative work, making it quite time-consuming and labor-intensive.

Apart from attendance tracking, lecturers are also required to provide and assess multiple assignments, with the number of assessments varying based on the subject's nature. Subjects with a higher weightage for summative assessments (coursework) compared to formative assessments (final exams) demand even more evaluations. The COVID-19 pandemic has further complicated this process as all classes have shifted to an online format. Consequently, students now submit their assessments digitally, with lecturers sharing submission links via platforms like WhatsApp and Telegram or send directly to the cloud storage. However, this online submission process has introduced new challenges. Many students failed to adhere to submission guidelines, which often use inconsistent filenames, varying file formats, and different file sizes. As a result, lecturers must manually rename, organize, and sometimes repair these files, which becomes an exhaustive task, especially when managing a large number of students and multiple online assessments.

A study by Panisoara (2020) highlighted the obstacles faced by lecturers during the COVID-19 pandemic, showing that numerous experienced significant stress due to unclear job roles, difficulties in managing online classes, and the fatigue associated with prolonged online teaching. Furthermore, the students became uninterested and lacked responses through online teaching. Without the use of appropriate tools, students may submit their work after the deadline, or in some instances, resort to plagiarism or copying from others through shared cloud storage accessible to all, allowing unethical downloading, editing, and misuse.

To address these challenges, this article introduces an innovative solution: A comprehensive web-based application integrating attendance tracking and assessment submission. This online system would allow the learners to log in to their attendance through mobile phones and upload the softcopy of assessments effortlessly, ensuring consistency in file formats and preventing file corruption. As a result, it would reduce lecturers' workload, with most clerical tasks being automated by the system.

This paper is organized as follows: It begins with a brief introduction to the project, followed by an exploration of related research and previous works in this field. This study outlines the research methodology before moving on to an analysis and discussion of the project's findings. Lastly, suggestions for potential improvement or recommendation of the project are presented in the Conclusion section.

2. LITERATURE REVIEW

As countries grappled with the challenges brought on by the COVID-19 pandemic, a wide range of teaching and learning methodologies have been introduced to sustain the continuation of the learning process and ensure that educational systems continue to function effectively. This unprecedented situation forced a rapid transition from traditional classroom teaching to digital or virtual learning settings.

In response to the pandemic, most educational institutions, from schools to universities, were compelled to temporarily close their physical campuses and shift their classes and lectures entirely online. This sudden and unplanned transition from offline to online learning posed numerous challenges for educators, students, and even parents. Many institutions were unprepared for such a drastic shift, resulting in a steep learning curve as they adjusted to new technologies and online platforms. As highlighted by Gillett-Swan (2017), the crucial problem was the lack of knowledge and skills in online teaching and learning, which led to disruptions in lesson planning, delivery, and student engagement. Another significant challenge of this pandemic was the observation of learners' performance through the online learning context. According to Mohammed et. al (2020), evaluating students' abilities, particularly in areas that require practical skills, technical competencies, and hands-on experience, became increasingly difficult in an online setting. For instance, courses that rely heavily on lab work, field studies, or teaching practicum faced hurdles in replicating these experiences in a virtual environment. In addition, students learning such subjects found it challenging to gain the same level of understanding and skill development as they would in a traditional classroom or laboratory setting.

Moreover, teachers struggled to design assessments that could accurately measure students' practical knowledge and competencies in an online format. This situation raised concerns about academic integrity, as the remote nature of assessments made it easier for some students to be involved in unethical practices. Consequently, teachers had to explore other approaches for assessment types, such as open-book examinations, group project assignments, and online presentations, to ensure that the learning outcomes of students are fairly evaluated. Furthermore, teachers were required to manually monitor and record student attendance, as well as calculate the total absenteeism percentage for each semester. This manual process posed several challenges, including the risk of buddy signing, misplaced attendance sheets, and difficulties in controlling student absenteeism (Zuanuwar, 2020).

In response to the evolving demands of teaching methods brought about by the global pandemic, educational experts have developed innovative online applications designed to manage various aspects of student learning more effectively. These advanced systems are capable of handling multiple tasks, such as tracking student attendance, facilitating assignment submissions, and streamlining the evaluation process. Many higher education institutions have adopted various Student Attendance Management Systems encompassing internet-based solutions such as web-based and mobile-based systems, as well as other computerized attendance methods. As highlighted by Anitha et al. (2016), a web-based Attendance Management System incorporating SMS technology to inform parents about their children's attendance significantly boosts students' motivation and sense of responsibility in attending classes. This proactive approach ensures that students are more likely to be present and accountable.

Furthermore, some universities have implemented QR code technology to record attendance, which allows attendance data to be stored directly on the server (Anitha et al., 2016). The use of QR codes accelerates data entry and minimizes errors, making the process more efficient and reliable. In many higher education institutions, lecturers evaluate overall attendance throughout the semester to determine if students are eligible to sit for their final examinations (Benyo et al., 2012). This emphasizes the importance of accurate attendance records and how technology can streamline this critical aspect of academic life. Vadvala (2024) proposed a digital attendance system to minimize the time spent on traditional pen-and-paper methods. His innovative solution integrates Radio Frequency Identification (RFID) sensors, Bluetooth devices, Wi-Fi networks, facial recognition, and fingerprint-based decentralized sensors, resulting in a highly accurate and nearly error-free attendance tracking process. Ishaq and Bibi (2023) have proposed similar project to build a real-time attendance tracking system by combining the RFID, the Google Sheets forms and the Internet of Things (IoT), to ensure accuracy of attendance monitoring.

In addition to attendance management, universities have increasingly encouraged the use of Learning Management Systems (LMS) for submitting assessments. As an example, the University of Southern Queensland (USQ) in Australia introduced two LMS platforms, Writely and Moodle, which enabled

students to send their assignments online (Petrus & Sankey, 2007). Both learners and educators replied encouragingly to these platforms, as they simplify the assignment submission process and facilitate timely feedback. The convenience and efficiency provided by these systems enhance the overall teaching and learning experience.

Almost the same application system has been constructed focusing on the submission of assignments, online grading, and viewing of assessment results (Sam, 1998). This system allows lecturers to upload files or question papers, grant authorized students access to view and respond to questions, export grades to a spreadsheet, send a warning message about assignment deadlines to students, and share graded assessments with students. Such systems not only streamline the assignment submission process but also ensure that students and lecturers can manage coursework more effectively. Abu Bakar (2025) asserts that online submission systems positively influence educational quality, promote learner autonomy, and improve institutional efficiency. Additionally, they enhance grading transparency, enable timely feedback, and support the development of digital literacy among both students and educators.

Moreover, Eaganathan and Maruf (2018) proposed an advanced assessment submission system, which incorporates encryption algorithms and cryptographic technology to secure assignment submissions, ensuring data integrity and confidentiality. One of the notable features of this system is that students can check the plagiarism percentage before submitting their assignments, which encourages originality and reduces academic misconduct. After the lecturer has graded the assignment, students are then allowed to view their work and the associated feedback. The system also sends automated notifications via email or SMS to remind students of upcoming deadlines or late submissions. These features significantly enhance the system's integrity and improve overall user satisfaction.

Research in tertiary education has demonstrated that the implementation of Learning Management Systems (LMS) significantly enhances the assessment experience. These systems contribute to a marked reduction in paper consumption (Ellis & Reynolds, 2013) and simplify the management of grading (Ellis & Reynolds, 2013). Additionally, LMS provides more space for comments and enables marking from any location (University of Glamorgan, 2012), besides incorporating rubrics and in-text comments (Ellis & Reynolds, 2013), as well as allowing for quicker grading of certain assessments (Buckley & Cowap, 2013). Overall, these advancements facilitate a more efficient and effective approach to student assessment and feedback in higher education.

3. METHODOLOGY

The Integrated Web-Based Record Management System for monitoring students' attendance and assignment submissions was designed using the System Development Life Cycle (SDLC) waterfall model. According to Dora and Dubey (2013), the SDLC is a structured approach consisting of five distinct phases: Analysis, Design, Implementation, Testing, and Maintenance. The process begins with a comprehensive analysis to understand system requirements, followed by designing the system architecture. Implementation involves the real codes, while the testing phase ensures that the system fulfills the user requirements and reliability. Finally, the maintenance phase addresses any system updates or issues, ensuring the system remains efficient and effective over time. This structured approach secures the system's robustness and adaptability.

Fig. 1 illustrates the flow of the System Development Life Cycle (SDLC) phases, progressing sequentially from the analysis stage to the enhancement or maintenance phase. The SDLC framework offers developers the flexibility to revisit earlier phases, allowing for necessary modifications or corrective actions when required. The analysis phase is particularly critical within the SDLC process, as it involves identifying the core business requirements to guide the project. During this phase, data collection is essential, and the gathered information will be verified, frequently revised, and tallied with the original system functionality to ensure accuracy and relevancy. Defining the project's problem statements is imperative, as it ensures

that the objectives are measurable and achievable. To support this process, formal meetings with real users should be engaged as the key elements for data collection, providing valuable insights into user needs and expectations. This comprehensive approach confirms that the system development aligns with organizational goals and user requirements.

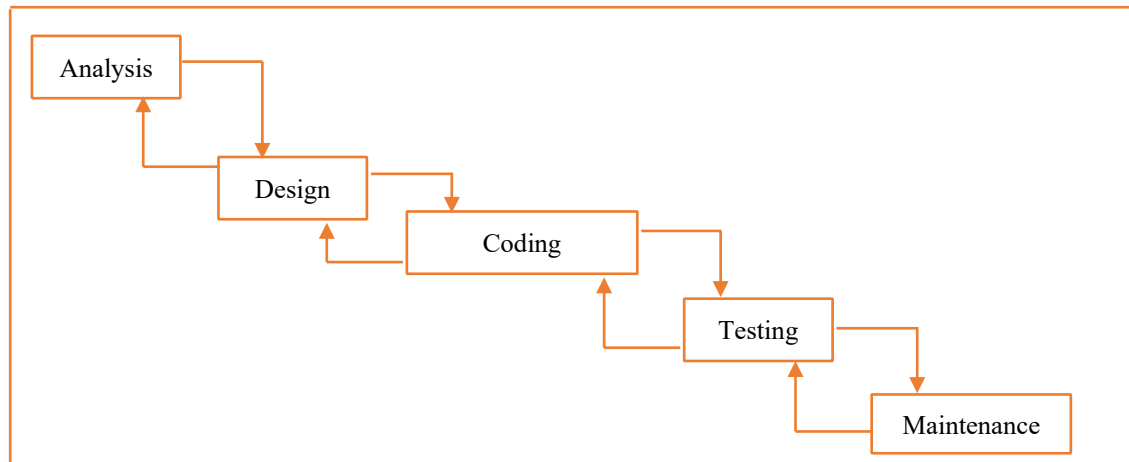


Fig. 1. System Development Life Cycle (SDLC)

The analysis documentation was presented, approved, and validated by primary users and senior management. Following this, the design phase commenced, focusing on incorporating an integrated database with a strong integrity structure, system interactivity, functionality, and relevant aspects of current technology demands. To expedite the development process, the prototyping methodology was employed, which allowed the design, implementation, and testing phases of the SDLC to be executed concurrently (Tavolato & Vincena, 1984). This approach significantly reduced project duration and costs. The initial functional prototype was tested by primary users and developers, and any necessary corrective actions were promptly undertaken to enhance the system. It was iteratively refined and re-evaluated until no further issues were identified by the users. In addition, the developers carried out thorough testing, taking into account factors like database integrity, network operation, concurrency issues, and security.

Following several SDLC cycles, the system was successfully deployed and is now being used by degree and diploma students at the Faculty of Computer & Mathematical Sciences, Universiti Teknologi MARA, Perlis branch. It is now in the process of being expanded to all students at the faculty level, demonstrating the system's scalability and adaptability. This thorough and structured approach to testing and system evaluation ensures the development of a reliable and efficient web-based system.

The system was evaluated by the same cohort of students to assess its functionality. Approximately 111 students have responded to the online questionnaire. An online questionnaire was designed, requiring students to provide sincere feedback on various aspects, including the system's design, features, requirements, and overall functionality. The questionnaire utilized a 5-point Likert scale, ranging from "Strongly Disagree" to "Strongly Agree," to capture the students' perceptions and opinions on the system's effectiveness. To ensure the questionnaire's reliability, the items' internal consistency was evaluated using Cronbach's alpha. This coefficient assesses how well the items in a scale are related to one another and whether they collectively measure the intended construct. With values ranging from 0 to 1, a higher Cronbach's alpha indicates stronger reliability. For this study, Cronbach's alpha was calculated for the 12 items, confirming that the instrument was consistent and reliable for measuring the intended variables.

The responses were analyzed using descriptive statistics, including the calculation of mean scores and standard deviations for each questionnaire item to determine overall trends in student feedback. The means were examined to provide deeper insights into the data, helping to identify specific strengths and areas for improvement. The analysis also included hypothesis testing using a one-sample t-test to determine whether the mean scores for all 12 questions were significantly greater than 4. This approach was used to confirm that the average scores exceeded the benchmark value, reflecting positive feedback from the students. It aims to gather insights that will guide future enhancements. By analyzing the students' responses, the development team will identify areas for improvement, ensuring that the system effectively meets user needs and expectations. This feedback loop is crucial for fostering continual growth and refinement in system performance.

4. RESULT AND DISCUSSION

The Integrated Web-Based Record Management System for Students' Attendance and Assessment Submissions was developed using HTML and PHP, with MySQL serving as the database. As outlined in the previous section, this web-based system comprises two main modules: The first records student attendance, while the second handles the management of students' assessment submissions. The following table outlines the business requirements for both modules, as specified by the users.

Table 1. Points as recommended by the users for the Attendance and Assessment Submission module

Attendance Module	Assessment Submission Module
- Students are permitted to log their attendance only during the class. - Students can upload documentation as evidence if they are absent for a valid reason. - Lecturers can update attendance if students forget to log in during the class. - Students can view their attendance records. - The system displays the attendance status and the percentage of attendance up to the current date. - The system can send email notifications to students should they be absent. - The system generates a list of absentees for a specific class and date. - Only authorized users are permitted to log in to the web-based system. - Only documents in valid formats and within the specified file size limit are allowed to be uploaded. - Students can update their information, such as contact numbers and electronic mail. - The system provides a detailed analysis of students' attendance records, reflecting their overall attendance performance.	- Only authorized users are permitted to upload assessments. - Students select the type of assessment they wish to upload. - Only documents in valid formats and within the specified file size limit are allowed to be uploaded. - Students can verify and confirm their assignment submission in the system. - Lecturers can easily download submitted assessments from the system before grading. - Graded assessments can be easily re-uploaded to the system before being returned to the students. - Students can view their assessment grades and download the graded assignments. - Students can also view their overall marks or coursework grades.

Fig. 2 displays the front page of the Attendance & Assessment Management Systems or AAMS.

Fig. 2. The main menu of AAMS

The following page has two main modules, namely attendance and assessment submissions. Access for the users is differentiated by the colour of the buttons; the blue button is specifically for students while the pink buttons are for the educators and the system administrator. The pink buttons are hidden from students, as they are restricted by the IP (Internet Protocol) address of the system administrator. Basic student profiles, including student matric number, student name, program, part, and course code, are retrieved from the Student Information Management Systems (SIMS) and integrated into the AAMS database, ensuring that only authorized students can access the system. To log in for attendance, students must enter their matric number and subject code. Once logged in, the system records a timestamp of their attendance. Access is permitted only during the allocated class time or designated slot.

Meanwhile, Fig. 3 illustrates a platform for students to upload documents, such as the medical certificate form, or any relevant documents as evidence if they cannot attend their class. The student must upload the documents or evidence and make sure the file is in PDF format. The system only accepts the size of a document of less than 2 MB. The system will upload the documents for the authorized student ID only. This feature is applied for assessment submission and declaration of absenteeism modules.

Fig. 3. Interface to upload the evidence of absenteeism

Fig. 4 provides a special feature in which the users can view details of individual attendance records and assessment submissions. From this figure, the students can keep track of their details record of attendance and proof that assignments have been submitted successfully.

ATTENDANCE & ASSESSMENT MANAGEMENT SYSTEMS (AAMS)

Nombor Matrik : 2020626504
 Nama Pelajar : ALIN AFINA BINTI JEMSUL
 Kod Kursus : ITS232
 Kumpulan : RCS1434E
 No. Phone : 0193737092
 Email Address : alinafinaaa@gmail.com

REKOD KEHADIRAN

#	Tarikh	Hari	Masa	Catatan
1	30/03/2022	RABU	14:54:02 PM	
2	06/04/2022	RABU	14:03:59 PM	
3	13/04/2022	RABU	14:06:07 PM	
4	20/04/2022	RABU	14:02:04 PM	
5	27/04/2022	RABU	14:02:01 PM	

Peratus Kehadiran Keseluruhan : 17.86%

REKOD PENYERAHAN ASSESSMENT

#	Masa Penghantaran	Assessment	Nama File	Markah
Tiada rekod dijumpai				

Fig. 4. Interface of students attendance and assessment submission records

Additionally, on this page of the platform, the students will be notified about their attendance performance, and they have to make sure that the percentage of class attendance is more than 80% in 14 weeks of the academic calendar. If they fail to fulfil this regulation, they cannot sit the final examination as the penalty. The students can check their assignment results; assessments that have been marked by the lecturer will be returned through this platform. This interface serves as a one-stop page, allowing students to monitor their attendance and coursework progress.

Administrators or lecturers have special privileges to review the formative assessment marks and the current attendance percentage, as illustrated in Fig. 5 and Fig. 6 below.

UNIVERSITI TEKNOLOGI MARA CAWANGAN PULAU PINANG, KAMPUS PERMATANG PAUH									
ATTENDANCE & ASSESSMENT MANAGEMENT SYSTEMS (AAMS)									
Rekod Kehadiran Terperinci									
#	No Matrik	Nama Pelajar	Kod Subjek	Kumpulan	Tarikh Kehadiran	Jumlah Kehadiran	Peratusan Kehadiran	Dokumen MCE/EL	Catatan
1	2020497986	AFIFAH ZUHAIKRAH BINTI ZUHAIMI (0192844809)	ITS232	RCS1434A	28/03/2022, 29/03/2022, 04/04/2022, 05/04/2022, 11/04/2022, 12/04/2022, 18/04/2022, 25/04/2022, 26/04/2022,	9	32.14%		BAR
2	2020469306	AFZA IRDINA BINTI YUSOF (0182057696)	ITS232	RCS1434A	28/03/2022, 29/03/2022, 04/04/2022, 05/04/2022, 11/04/2022, 12/04/2022, 18/04/2022, 25/04/2022, 26/04/2022,	9	32.14%		BAR
3	2020625704	AHMAD HASIF HAIKAL BIN MOHD SUHAIMI (0193704623)	ITS232	RCS1434A	28/03/2022, 29/03/2022, 04/04/2022, 05/04/2022, 11/04/2022, 12/04/2022, 18/04/2022, 25/04/2022, 26/04/2022,	9	32.14%		BAR

Fig. 5. Interface of student attendance performance

UNIVERSITI TEKNOLOGI MARA CAWANGAN PULAU PINANG, KAMPUS PERMATANG PAUH							
ATTENDANCE & ASSESSMENT MANAGEMENT SYSTEMS (AAMS) COURSEWORK DETAILS							
SENARAI MARKAH KERJA KURSUS							
#	No Matrik	Nama Pelajar	Kod Subjek	Kumpulan	Assessment	Markah	Jumlah Markah
1	2020497986	AFIFAH ZUHAIKRAH BINTI ZUHAIMI (0192844809)	ITS232	RCS1434A			0
2	2020469306	AFZA IRDINA BINTI YUSOF (0182057696)	ITS232	RCS1434A			0
3	2020625704	AHMAD HASIF HAIKAL BIN MOHD SUHAIMI (0193704623)	ITS232	RCS1434A			0
4	2020876892	AININ SABIHA BINTI MOHD SAYUTHI (019-7914249)	ITS232	RCS1434A			0
5	2020604464	AQILAH NUWAIKRAH BINTI AMIZUL ANUAR (01161708226)	ITS232	RCS1434A			0

Fig. 6. Interface of student formative assessment marks (all marks are 0 as the marking of assessment submitted is still ongoing)

The feedback on system functionalities was given and collected from students that already used the system for one semester. An online questionnaire was constructed and shared with the students to give feedback and views for future system enhancements. The internal consistency of the survey instrument was assessed using Cronbach's alpha coefficient, which evaluates how reliably the items in the questionnaire measure the same underlying construct. The Cronbach's alpha coefficient for the 12-question survey was found to be 0.965, indicating a high level of internal consistency, as values above 0.9 are generally considered excellent (George & Mallery, 2024).

After evaluating the questionnaire's reliability, the analysis proceeded with descriptive statistics to summarize and interpret the data, providing an overview of the respondents' demographic information and their feedback on the system. Table 2 displays the students' demographic information, offering insights into the sample's composition.

Table 2. Demographic

Variable	Category	Frequency	Percentage
Level	Diploma	41	36.9
	Degree	70	63.1
Gender	Male	24	21.6
	Female	87	78.4
Semester	Sem 1	104	93.7
	Sem 2	7	6.3

The demographic analysis of the survey participants provides a clear overview of the distribution of respondents based on academic level, gender, and semester. A total of 111 students participated in this study, with the majority being degree students, accounting for 63.1% ($n = 70$), while diploma students made up 36.9% ($n = 41$). This indicates that the system was evaluated by a diverse group of students, though degree students represented a larger proportion of the sample. In terms of gender, the distribution showed a notable majority of female respondents, comprising 78.4% ($n = 87$), compared to male respondents, who accounted for only 21.6% ($n = 24$). The gender imbalance may suggest a stronger representation of female students in the surveyed cohort, potentially reflecting the demographics of the academic programs involved or the accessibility of the survey.

The majority of respondents were first-semester students, representing 93.7% ($n = 104$), while second-semester students constituted a much smaller group, at only 6.3% ($n = 7$). It suggests that the system was predominantly evaluated by students new to their academic programs. First-semester students are likely to engage more frequently with tools like the AAMS system as they adjust to academic expectations, which might explain their higher representation.

Following the demographic analysis, this study proceeded to examine the mean scores and standard deviations for each of the 12 questions in the questionnaire. This analysis aimed to provide a clearer understanding of the central tendencies and variability in the respondents' feedback. The results are summarized in Table 3, which presents the mean and standard deviation values for all 12 items.

Table 3. Mean and standard deviation of feedback on survey questions

Item	Mean	Std Dev
Q1 - I find the AAMS system very useful.	4.6036	0.5765
Q2 - I find the interface of AAMS easy to use.	4.5315	0.64413
Q3 - I have a positive attitude on using AAMS.	4.5495	0.61406
Q4 - Overall, I like to use AAMS.	4.5045	0.64489
Q5 - I will recommend all lecturers to apply AAMS.	4.4865	0.68576
Q6 - I find the system running smoothly without failures.	4.3243	0.67638
Q7 - I am satisfied with the system's interface.	4.4324	0.64133
Q8 - I am satisfied with the system's features.	4.4505	0.5837
Q9 - Features provided in AAMS are really needed by all students.	4.4414	0.70948
Q10 - The output/reports from AAMS are really valuable.	4.5135	0.60098
Q11 - The system is accessible anywhere.	4.5586	0.59825
Q12 - The system is accessible anytime.	4.5676	0.66908

Following the presentation of the table, the descriptive statistics for the AAMS system survey painted a compelling picture of students' perceptions and experiences with the system. Specifically, the mean and standard deviation values for questions Q1 to Q12 highlighted consistent positive feedback. The mean scores for all questions exceeded 4.3, underscoring strong agreement among respondents regarding the

system's effectiveness. Notably, Q1 ("I find the AAMS system very useful") recorded the highest mean of 4.6036, signifying widespread recognition of the system's utility. Meanwhile, Q6 ("I find the system running smoothly without failures") recorded the lowest mean of 4.3243, though it remained firmly within the positive range. These results collectively suggest that students view the AAMS system as an essential tool that simplifies critical academic processes, such as attendance tracking, assignment submission, and evaluation.

In terms of variability, the standard deviation values ranged between 0.5765 and 0.70948, showing a relatively low to moderate spread in students' responses. For instance, Q1 has the lowest standard deviation of 0.5765, signifying a high level of agreement and consistency among students about the usefulness of the system. Similarly, Q11 ("The system is accessible anywhere") and Q10 ("The output/reports from AAMS are really valuable") exhibited low standard deviations of 0.59825 and 0.60098, respectively, suggesting a strong consensus regarding the accessibility and value of the system's outputs. On the other hand, Q9 ("Features provided in AAMS are really needed by all students") has the highest standard deviation of 0.70948, indicating a slightly broader range of opinions among students about the necessity of the system's features.

Comparing the mean and standard deviation values collectively, it is evident that questions with higher mean scores tend to generally have lower standard deviations, reflecting a strong alignment in students' positive views. For example, Q1 and Q11, which both have high mean scores, also displayed lower standard deviations, reinforcing the consistency in students' agreement. Conversely, questions with slightly lower mean scores, such as Q6 and Q9, demonstrated relatively higher standard deviations, suggesting that these aspects of the system may be perceived differently by a small proportion of students. This disparity might point to areas for further exploration or enhancement to meet the diverse expectations of users.

Overall, the combination of high mean scores and low to moderate standard deviations across all questions indicates that the AAMS system is well-received by students, with minimal disagreement. However, a closer examination of questions with higher variability, including Q9, could provide valuable insights into potential areas for improvement. These findings support the conclusion that the system has successfully met its objectives, while also highlighting opportunities for refining its features to cater to broader user needs.

To further validate the statistical significance of the findings, a one-sample t-test was conducted to determine whether the mean scores for the survey items were significantly higher than the benchmark value of 4, which represents agreement on the Likert scale. This test assesses whether students' positive perceptions of the AAMS system are statistically significant and exceed the threshold for general agreement. The results of the t-test, as presented in Table 4, confirmed that all mean scores for Q1 to Q12 were significantly greater than 4, with p-values less than 0.05. These results provide strong evidence that the AAMS system is not only well-received but also highly effective in fulfilling its purpose of streamlining attendance tracking, assignment submission, and evaluation processes. High mean scores, such as ≥ 4.0 on a 5-point Likert scale, usually reflect positive perceptions or satisfaction and show that respondents agree or strongly agree with the statements. The interpretation of mean values above the neutral midpoint as high/positive perception is supported by research literature (Rokeman, 2024).

Table 4. One-Sample T-Test Results (Test value > 4)

Item	p-value (1- tailed)
Q1 - I find the AAMS system very useful.	0.000
Q2 - I find the interface of AAMS easy to use.	0.000
Q3 - I have a positive attitude on using AAMS.	0.000
Q4 - Overall, I like to use AAMS.	0.000
Q5 - I will recommend all lecturers to apply AAMS.	0.000
Q6 - I find the system running smoothly without failures.	0.000
Q7 - I am satisfied with the system's interface.	0.000
Q8 - I am satisfied with the system's features.	0.000
Q9 - Features provided in AAMS are really needed by all students.	0.000
Q10 - The output/reports from AAMS are really valuable.	0.000
Q11 - The system is accessible anywhere.	0.000
Q12 - The system is accessible anytime.	0.000

It can be confidently concluded that students' perceptions of the AAMS system are statistically significant and highly positive. These findings demonstrate that the AAMS system is not only perceived as useful and efficient by most students but also significantly meets their expectations in terms of functionality, ease of use, and accessibility. The consistently high ratings across all survey items suggest that the system has effectively achieved its objectives in streamlining attendance tracking, assignment submission, and evaluation processes. Moving forward, these positive results will provide valuable insights into further refinement and potential widespread adoption of the system.

5. CONCLUSION

In conclusion, the integration of technology in attendance management and assessment submission systems has proven to be highly effective in higher education institutions. These advanced systems not only facilitate efficient data management and reduce errors but also enhance communication, accountability, and the overall educational experience for students and lecturers alike. The positive feedback from students regarding the AAMS system demonstrated by the high mean scores and statistically significant results highlights its effectiveness in meeting its intended objectives. The system has significantly streamlined key academic processes, such as attendance tracking, assignment submission, and evaluation, ultimately contributing to a more efficient and organized learning environment.

6. ACKNOWLEDGEMENTS/FUNDING

The authors would like to acknowledge the support of Universiti Teknologi Mara (UiTM), Cawangan Negeri Pulau Pinang, Kampus Permatang Pauh, Malaysia for providing the data and facilities on this research.

7. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

8. AUTHORS' CONTRIBUTIONS

Jamal Othman: Introduction, Literature Review & Methodology; **Siti Balqis Mahlan:** Data Analysis, Result & Discussion; **Maisurah Shamsuddin:** Result & Discussion; **Syarifah Adilah Mohamed Yusoff:** Literature Review & Editing; **Zalina Othman:** Data Collection & System Development.

9. REFERENCES

- Abu Bakar, S. (2025). The impact of online submission systems on teaching practice in digital education environments: A case study of Open University Malaysia. *Muallim Journal of Social Sciences and Humanities*, 9(S1), 231-239. <https://doi.org/10.33306/mjssh/345>
- AGCAS. (July, 2021). *Higher education lecturer*. Prospects. <https://www.prospects.ac.uk/job-profiles/higher-education-lecturer>
- Anitha V. P., Krishna A, Kshama P.M., Correa M. (2016). Web service for student attendance management system, 5(3).
- Benyo B, Sodor B, Doktor T, Fördös G. (2012). Student attendance monitoring at the university using NFC. In *Wireless Telecommunications Symposium 2012* (pp. 1-5). IEEE. <https://doi.org/10.1109/WTS.2012.6266137>
- Buckley, E. & Cowap, L. (2013). An evaluation of the use of Turnitin for electronic submission and marking and as a formative feedback tool from an educator's perspective. *British Journal of Educational Technology*, 44(4), 562-570. <https://doi.org/10.1111/bjet.12054>
- Dora, S.K. and Dubey, P. (2013). Software Development Life Cycle (SDLC) analytical comparison and survey on traditional and agile methodology. *Journal of Research in Science & Technology*, 2(8), 22-30.
- Eaganathan, U., & Md. Maruf Hasan, S. (2018). Proposed literature review on ASIA Pacific University assignment submission and feedback system for future development. *International Journal of Engineering & Technology*, 7(2.33), 480-483. <https://doi.org/10.14419/ijet.v7i2.33.14815>
- Ellis, C. & Reynolds, C. (2013). *EBEAM Final Report*. https://ipark.hud.ac.uk/sites/default/files/EBEAM_Project_report_compressed.pdf
- George, D., & Mallery, P. (2024). *IBM SPSS statistics 29 step by step: A simple guide and reference*. Routledge. <https://doi.org/10.4324/9781032622156>
- Gillett-Swan, J. (2017). The challenges of online learning: Supporting and engaging the isolated learner. *Journal of Learning Design*, 10(1), 20-30. <https://doi.org/10.5204/jld.v9i3.293>
- Ishaq, K. & Bibi, S. (2023). IoT based smart attendance system using RFID: A systematic literature review. *arXiv*. <https://doi.org/10.48550/arXiv.2308.02591>
- Mohammed, A. O., Khidhir, B. A., Nazeer, A., & Vijayan, V. J. (2020). Emergency remote teaching during the Coronavirus pandemic: The current trend and future directives at Middle East College Oman. *Innovative Infrastructure Solutions*, 5(3), 72. <https://doi.org/10.1007/s41062-020-00326-7>
- Panisoara, I.O.; Lazar, I.; Panisoara, G.; Chirca, R.; Ursu, A. S. (2020). Motivation and continuance intention towards online instruction among teachers during the COVID-19 pandemic: The mediating effect of burnout and technostress. *International Journal of Environmental Research and Public Health*, 17(21), 8002. <https://doi.org/10.3390/ijerph17218002>

- Petrus, K., & Sankey, M. (2007, January). Comparing Writely and Moodle online assignment submission and assessment. In *Proceedings of the 3rd International Conference on Pedagogies and Learning*. <https://www.researchgate.net/publication/242044743>
- Rokeman, N. R. M. (2024). Likert measurement scale in education and social sciences: Explored and explained. *EDUCATUM Journal of Social Sciences*, 10(1), 77-88. <https://doi.org/10.37134/ejoss.vol10.1.7.2024>
- Sam H. (1998). HWSAM: A web-based automated homework submission system. In *1998 ASEE/IEEE Frontiers in Education Conference Proceedings (FIE '98)* (pp. 580-582). IEEE. <https://doi.org/10.1109/FIE.1998.738744>
- Sanusi, N., Zulkifli, H., & Hamzah, M. I. (2025). Digital technology in classroom assessment: A bibliometric study. *TEM Journal*, 14(1), 551-561. <https://doi.org/10.18421/TEM141-49>
- Tavolato, P., & Vincena, K. (1984). A prototyping methodology and its tool. In *Approaches to Prototyping* (pp. 434-446). Springer. https://doi.org/10.1007/978-3-642-69796-8_38
- University of Glamorgan (2012). *Evaluation of assessment diaries and grade mark at the University of Glamorgan*. http://jiscdesignstudio.pbworks.com/w/file/67927168/JISC_Assessment_and_Feedback_Glamorgan_Final_Project_Report_post_JISC_feedback_Nov_2012.docx
- Vadvala, A. (2024). A comprehensive review of digitalized attendance systems for various institutions. *International Journal of Engineering Applied Sciences and Technology*, 09(03), 119-123. <https://doi.org/10.33564/IJEAST.2024.v09i03.011>
- Zuanuwar, M. (2020). The influence student attendance management system on academic performance. *Journal of Social Transformation and Education*, 1(1), 26-62.



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Hydrological Risk Assessment of the Johor River: Flood Frequency Analysis with L-Moments

Nur Hanida Othman¹, Basri Badyalina^{2*}, Sheril Haiza Shah Izan³, Ani Shabri⁴,
Muhammad Fadhil Marsani⁵, Fatin Farazh Ya'acob⁶

^{1,3}Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Negeri Sembilan Branch, Seremban Campus, 70300 Negeri Sembilan, Malaysia.

²Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Johor Branch, Segamat Campus, 85000 Johor, Malaysia.

⁴Department of Mathematical Sciences, Faculty of Science, Universiti Teknologi Malaysia, 81310, Johor, Malaysia.

⁵School of Mathematical Sciences, Universiti Sains Malaysia, 11800 Minden, Penang, Malaysia.

⁶Faculty of Business and Management, Universiti Teknologi MARA Johor Branch, Segamat Campus, 85000 Johor, Malaysia.

ARTICLE INFO

Article history:

Received 29 April 2025

Revised 4 July 2025

Accepted 16 August 2025

Published 1 September 2025

Keywords:

L-Moments
Generalized Pareto
Return Period
Hydrological Modelling
Flood Frequency Analysis
Johor River

DOI:

10.24191/jcrinn.v10i2.527

ABSTRACT

Flooding remains one of the most significant hydrological hazards in Malaysia, particularly within the Johor River Basin. This study conducts a Flood Frequency Analysis (FFA) using the L-Moment method to identify the most suitable probability distribution for modelling annual maximum streamflow at the Kahang River station. A 45-year dataset (1978-2022) was analyzed using three candidate distributions: Generalized Logistic (GLO), Generalized Pareto (GPA), and Pearson Type-III (PE3). The L-Moment Ratio Diagram (LMRD), alongside statistical performance metrics such as *MAE*, *RMSE*, *MAPE*, *RMSPE*, *R*² and Euclidean Distance was employed to evaluate model accuracy. Results reveal that the GPA distribution provides the best fit as it provides smallest value of *MAE*, *MAPE*, *RMSPE* and Euclidean Distance, demonstrating superior accuracy in predicting extreme flood events, particularly for high return periods. The study offers critical insights for flood risk assessment, infrastructure planning, and early warning system development in the region. Lastly, it provided researchers with flood analysis methods using L-moments and extreme value distributions.

1. INTRODUCTION

Flood is one of the most destructive natural disasters, with severe socio-economic impacts globally. In Malaysia, the Johor River basin is highly flood-prone due to its tropical climate, featuring heavy rainfall during the northeast monsoon (November-February). High temperatures (24°-32°C) and year-round humidity further increase flood risks, making Johor particularly vulnerable to extreme weather. Over past

^{2*} Corresponding author. E-mail address: basribdy@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.527>

20 years, major flood such as those in 2006-2007, 2011 and 202 have caused widespread displacement, damage to infrastructure and livelihood disruptions (Ahmad et al., 2020)

To assess flood risks, hydrologists use various Flood Frequency Analysis (FFA) models, the key method for studying extreme river flows and predicting recurrence intervals (Hamed & Rao, 2019; Jan et al., 2016). FFA methods evolved to incorporate historical data, censored peaks, and outlier tests, significantly improving flood quantile estimation (Stedinger, 1993; Hamed & Rao, 2019; Ali & Rahman, 2022). Early foundational work (Fakhru'l-Razi et al., 2020) by Todd (1957) and Linsley (1986) established streamflow analysis principles, while later Moughamian et al. (1987) advanced the distribution methods.

Advanced statistical techniques, such as L-moments presented an efficient alternative for the estimation of extreme hydrologic data (Badyalina et al., 2021). The reason to utilized the L-Moment in this study is because L-moments are less sensitive to outliers and produce better parameter estimation of probability distributions like the Generalized Logistic (GLO), Generalized Pareto (GPA) and Pearson Type-III (PE3) distributions (Marsani et al., 2022; Hassim et al., 2022). By applying these methods to historical records of floods in the Johor River, it was then feasible to enhance the description of the frequency and size of flood events and can refine the precision of flood risk predictions.

Studies in Johor rivers, such as the Sayong River, Segamat River, and Johor's Region III confirm GLO's superior fit for peak flow data (Zamani et al., 2024; Badyalina et al., 2021; Che Ilias et al., 2021). Meanwhile, GPA outperformed other distributions in the Segamat River (Romali & Yusop, 2017) and matched regional L-moments in flood-prone areas like Northern Iran (Adhami, 2024). The Pearson Type-III (P3) distribution handles asymmetric data through shape, scale, and location parameters, offering versatility for skewed flood data (Zhang et al., 2012). Recent studies highlight its superiority in the Johor River Basin (Jafry et al., 2023; Badyalina et al., 2022) and moderate-skew environments (Turhan, 2022), even outperforming GLO and GPA in some cases (Valentini et al., 2024). Though less common in Johor, PE3 shows promise for regional adaptation.

Hence, this study aims to improve flood risk assessment for Johor River by applying L-Moments and extreme value distributions (Generalized Pareto (GPA), Generalized Logistic (GLO) and Pearson Type-III (PE3)) to analyse historical flood data from Sungai Kahang station. The advantages of using each distribution are for GLO, it has satisfactory estimation of extreme flood events as it is able to capture tail behaviour well (Hamed & Rao, 2019). GPA is effective for heavy-tailed data meanwhile PE3 distribution can augment the data input from positively skewed data to negative skewed data depending on the location parameter, the PE3 distribution has great utility across many real-world scenarios.

The research will identify the most accurate distribution for modelling flood frequency, estimate return periods for extreme events, and provide actionable insights for flood management agencies like Department of Irrigation and Drainage Malaysia (DIDM). While the findings will enhance early warning systems and infrastructure planning, limitations include time constraints restricting long-term climate analysis and data accessibility challenges affecting prediction precision. The methodology used in this research is presented in methodology section, covering the study area description, probability distributions, and evaluation criteria for distributions performance.

2. METHODOLOGY

2.1 Study area

This study analyzes flood patterns and extreme streamflow events in Johor, Malaysia, using historical data from a river station which is Kahang River (Station ID 2235401, coordinates 2°17'42"N 103°34'39"E). The research utilizes a 45-year dataset from 1978 until 2022 of annual maximum streamflow (in m³/s) obtained from the Department of Irrigation and Drainage Malaysia (DIDM). This study focuses on annual

streamflow data, excluding weekly and monthly measurements from the analysis. These long-term discharge records, collected from the gauging station shown in Fig. 1., hydrological stations in Johor including Kahang River.

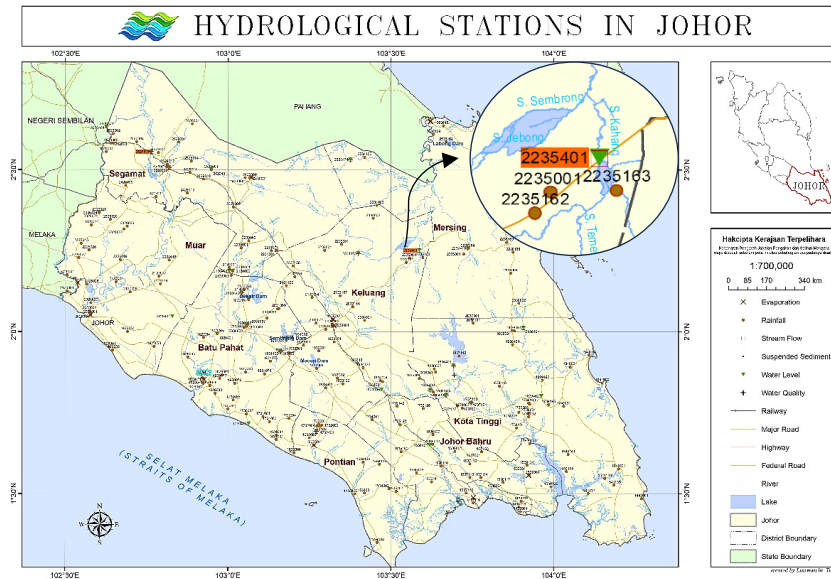


Fig. 1. Location of Kahang River

Source: Department of Irrigation and Drainage Malaysia

2.2 Gringorten plotting position

Gringorten's formula is commonly used to plot a set of ordered observations in a simpler format. According to Gringorten (1963), the Gringorten plotting position is used to find the optimal plotting position. This method is especially helpful in proving the efficiency of the plotting position, even when the sample size is less than 20. The Gringorten plotting position is calculated using the following formula:

$$P_i = \frac{i - 0.44}{n + 0.12} \quad (1)$$

where P_i is the plotting position for the i_{th} value, i is the rank of the value in the ordered dataset, and n is the sample size of the data.

2.3 L-Moments

L-moments (LMO), which are derived from probability-weighted moments (PWM), are commonly employed in hydrology for analysing extreme events like floods and droughts. They outperform traditional statistical methods by being less affected by outliers and working well with small datasets (Hamed & Rao, 2019). L-moments provide reliable estimates of distribution parameters, improving flood frequency analysis and water management decisions. Hosking (1990) developed the PWM approach that forms the basis for L-moment calculations. Assume $x_{1:n} \leq x_{2:n} \leq \dots \leq x_{n:n}$ represented the data in specific order with a sample size of n . Landwehr (1979) outlined the L-Moment method's unbiased sample estimator by the following formula:

$$b_r = \frac{1}{n} \binom{n-1}{r}^{-1} \sum_{i=r+1}^n \binom{i-1}{r} x_{i:n} \quad (2)$$

The first four elements of an unbiased sample estimator are listed below:

$$b_0 = \frac{1}{n} \sum_{i=1}^n x_{i:n} \quad (3)$$

$$b_1 = \frac{1}{n} \sum_{i=2}^n \frac{(i-1)}{(n-1)} x_{i:n} \quad (4)$$

$$b_2 = \frac{1}{n} \sum_{i=2}^n \frac{(i-1)(i-2)}{(n-1)(n-2)} x_{i:n} \quad (5)$$

$$b_3 = \frac{1}{n} \sum_{i=2}^n \frac{(i-1)(i-2)(i-3)}{(n-1)(n-2)(n-3)} x_{i:n} \quad (6)$$

Meanwhile, the first four sample estimates for L-Moments listed as:

$$l_1 = b_0 \quad (7)$$

$$l_2 = 2b_1 - b_0 \quad (8)$$

$$l_3 = 6b_2 - 6b_1 + b_0 \quad (9)$$

$$l_4 = 20b_3 - 30b_2 + 12b_1 + b_0 \quad (10)$$

Hence, the samples of L-Moment ratio are addressed as follows:

$$t_2 = \frac{l_2}{l_1} \quad (11)$$

$$t_3 = \frac{l_3}{l_2} \quad (12)$$

$$t_4 = \frac{l_4}{l_2} \quad (13)$$

L-moments summarize streamflow data, where l_1 (L-location) represents the central value, l_2 (L-scale) measures data variability (higher values indicate greater spread), and t_2 (L-CV) quantifies relative variation. L-skewness (t_3) and L-kurtosis (t_4) describe tail and peak behavior. Table 1 evaluates three distributions (GLO, GPA, PE3) for flood frequency analysis, showing their formulas ($x(F)$ = quantile at non-exceedance probability $F = 1 - 1/T$, where T = return period). Each distribution is defined by location (ξ), scale (α), and shape (k) parameters, which determine its fit to the data.

Table 1. Estimated distribution parameters using the L-Moments technique

Dist	Cumulative Density Function	Parameter Estimation
GLO	$x(F) = \hat{\xi} + \frac{\hat{\alpha}}{\hat{k}} \left(1 - \left(\frac{1-F}{F} \right)^{\hat{k}} \right)$	$\hat{k} = -t_3$
		$\hat{\alpha} = \frac{l_2}{\Gamma(\hat{k})[\Gamma(1-\hat{k}) - \Gamma(2-\hat{k})]}$
		$\hat{\xi} = l_1 - \frac{\hat{\alpha}}{\hat{k}} + \hat{\alpha}\Gamma(\hat{k})\Gamma(1-\hat{k})$
GPA	$x(F) = \hat{\xi} + \frac{\hat{\alpha}}{\hat{k}} \{1 - [1-F]^{\hat{k}}\}$	$\hat{k} = \frac{1-3t_3}{1+t_3}$
		$\hat{\alpha} = l_2(\hat{k}+1)(\hat{k}+2)$
		$\hat{\xi} = l_1 - \frac{\hat{\alpha}}{\hat{k}} + \frac{\hat{\alpha}}{\hat{k}(\hat{k}+1)}$
PE3	$x(F) = \hat{\xi}\hat{k} + K_T\sqrt{\hat{\xi}^2\hat{k}}$ where $K_T = \frac{2}{C_s} \left[\left\{ \frac{C_s}{6} \left(\Phi^{-1}[1-F] - \frac{C_s}{6} \right) + 1 \right\}^3 - 1 \right]$ $C_s = \frac{2}{\sqrt{\hat{k}}}$	$\hat{k} = 0.0127331632 + \frac{1.0246130369}{3p(t_3)^2} - \frac{1.0246130369}{3p(t_3)^2} -$ $\frac{0.0024863669}{(3p(t_3)^2)^2} + \frac{0.0001169073}{(3p(t_3)^2)^3} - \frac{0.0000027751}{(3p(t_3)^2)^4} + \frac{0.000000323}{(3p(t_3)^2)^5} -$ $\frac{0.000000001}{(3p(t_3)^2)^6}$
		$\hat{\alpha} = \frac{l_2}{2S_1(\hat{k}) - \hat{k}}$
		$S_1(\hat{k}) = \int_0^\infty \left[\int_0^x \frac{1}{\Gamma(\hat{k})} t^{\hat{k}-1} e^{-t} dt \right]^{\gamma} \frac{1}{\Gamma(\hat{k})} x^{\hat{k}} e^{-x} dx$ $\hat{\xi} = l_1 - \hat{k}\hat{\alpha}$

Source: Hamed and Rao (2019)

2.4 Accuracy Measure Performance

In this section, five accuracy performance measures are used which are Mean Absolute Error (*MAE*), Mean Absolute Percentage Error (*MAPE*), Root Mean Square Error (*RMSE*), Root Mean Square Percentage Error (*RMSPE*) and Coefficient of Determination (R^2). *MAE* is the mean of the absolute magnitude of differences between predicted and actual values irrespective of direction *MAPE* generalizes this by quantifying errors in percentage terms, thus it can be used to compare across different magnitudes of floods. *RMSE* conveys bigger errors by squaring the differences before computing the mean, thus it proves useful in identifying poorer predictions of the larger flood events. Lastly, *RMSPE* is a percentage version of *RMSE*, but is better suited to evaluate predictions that vary greatly in magnitude, such as those for a number of floods return periods. The R^2 determine how well the theoretical estimated value derived from distribution fits the actual data. The smallest error measure and larger value of R^2 implies that the data perform better for the model. Hence, the equation of each accuracy performance is presented in Eq. (14) through Eq. (18).

$$MAE = \frac{1}{n} \sum_{i=1}^n |F(y_i) - F(\hat{y}_i)| \quad (14)$$

$$MAPE = \frac{100}{n} \sum_{i=1}^n \left| \frac{F(y_i) - F(\hat{y}_i)}{F(y_i)} \right| \quad (15)$$

$$RMSE = \sqrt{\frac{\sum_{i=1}^n (F(y_i) - F(\hat{y}_i))^2}{n}} \quad (16)$$

$$RMSPE = \sqrt{\frac{1}{n} \sum_{i=1}^n \left(\frac{F(y_i) - F(\hat{y}_i)}{F(y_i)} \right)^2} \times 100 \quad (17)$$

$$R^2 = \frac{\sum_{i=1}^n (F(\hat{y}_i) - F(y_i))^2}{\sum_{i=1}^n (F(\hat{y}_i) - F(y_i))^2 + \sum_{i=1}^n (F(y_i) - F(\hat{y}_i))^2} \quad (18)$$

where n is the number of observations, $F(y_i)$ is the observed value, $F(\hat{y}_i)$ represented the predicted value, $\bar{y}$ is the mean of the observed value. The accuracy measurements such as MAE , $MAPE$, $RMSE$ and $RMSPE$ quantify the differences between the actual observation of annual peak flow and predicted estimates of peak flow from each distribution.

2.5 L-Moment ratio diagram

Hosking and Wallis (1997) introduced the L-Moment Ratio Diagram (LMRD) as a tool for calculating an accurate distribution of a catchment's streamflow series. Regarding a three-parameter distribution, the LMRD depicts the theoretical relationship between t_3 (skewness) and t_4 (kurtosis), as given in Eq. 12 and Eq. 13 respectively. The value of t_4 is calculated based on the Kahang River streamflow data and plot LMRD in order to recognize which distribution line lies closely.

$$t_4 = A_0 + A_1 t_3 + A_2 (t_3)^2 + A_3 (t_3)^3 + A_4 (t_3)^4 + A_5 (t_3)^5 + A_6 (t_3)^6 + A_7 (t_3)^7 + A_8 (t_3)^8 \quad (19)$$

The coefficients of A_k where $k=0, 1, 2, \dots, 8$ for the GLO, GPA and PE3 respectively are from the L-Moment. To measure the distance between observed and predicted values for each distribution, Euclidean Distance can be used. This is to strengthen the graphical evidence on which distribution is the best. Hence, the formula for Euclidean Distance (Krislock & Wolkowicz, 2012) is expressed by:

$$d(p, q) = \sqrt{\sum_{i=1}^n (q_i - p_i)^2} \quad (20)$$

where q_i and p_i is the Euclidean vectors, starting from the origin of the space (initial point). For Euclidean distance, the smallest distance from the baseline indicates better model performance.

3. RESULT AND DISCUSSION

Maximum streamflow also closely related with flooding engineering, flood management measurement development and flood system facilities drainage system. Table 1 summarized the descriptive statistics of Kahang River site's annual peak flow data from 1978 until 2022. It exhibits a mean peak flow of 280.41, higher than its median (193.95), indicating a right-skewed distribution with substantial variability. This large spread between the mean and median, coupled with a high standard deviation (262.633) with min (54.8) and max (1284.8), confirms the presence of extreme flow events in the dataset.

Table 2. Descriptive Statistics for Annual Streamflow of Kahang River

Station	n	Mean	SD	Min	Max	Skewness	Kurtosis
Kahang River	45	280.41	262.63	54.8	1284.8	2.00	3.87

The peak flow data for the Kahang River has a skewness of 2.0, suggesting that the data is skewed to the right with more frequent lower and occasional high peaks. The positive value of kurtosis suggests a leptokurtic (Kurtosis > 3.0), indicate that a higher likelihood of extreme flow events. Therefore, the results confirms that the peak flow data does not follow normal distribution but instead has heavy tail. To better capture this behavior, the study uses three non-normal probability distributions which are Generalized Logistic (GLO), Generalized Pareto (GPA) and Pearson Type-III (PE3) to model the annual maximum flows in Johor. Their parameters were estimated by applying the L-Moment, with the results listed in Table 3. To visualize the observed annual peak flow data, the Gringorton Plotting Position equation was used, as illustrated in Fig. 2.

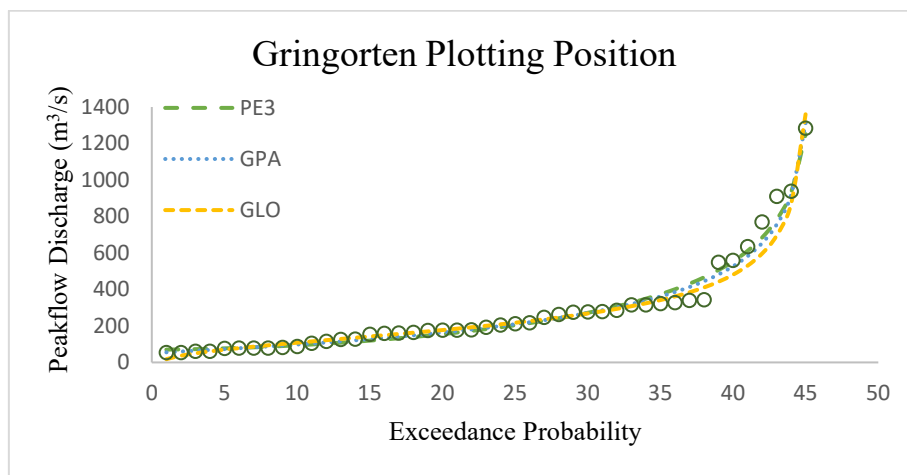


Fig. 2. CDF plot for the Kahang River annual peak flow and the candidate distributions

Fig. 2 shows all three distributions which are PE3, GLO and GPA which closely follow the observed data in the lower and central parts of the distribution, accurately capturing the general pattern of peak flows. However, in the upper tail region, where high discharge values are critical for flood risk assessments, there is a visible divergence among the models. While PE3 shows a slightly higher estimation, GPA tends to follow the observed data more consistently than GLO as the discharge increases. This suggests that GPA may be better at representing extreme events in this case, which is vital information when considering flood risk management and infrastructure design.

Table 3. Estimated distribution parameters using the L-Moment technique

Distributions	Parameters		
	ξ	α	$\hat{k}$
GLO	200.5983432	91.6763304	-0.4226313
GPA	53.2113263	184.416994	-0.1883087
PE3	280.412222	268.525419	2.553943

Table 3 displays the the estimated parameters for GLO, GPA, and PE3 distributions by employing the L-Moment technique. The GLO distribution has a central tendency (ξ) of 200.5983, the lowest variability (α) at 91.6763, and a negative skewness (k), indicating left-skewed data with more frequent low flood values. The GPA distribution has a much lower central tendency ($\xi = 53.2113$), higher variability ($\alpha = 184.4169$), and mild left-skewness ($k = -0.1883$). In contrast, the PE3 distribution has the highest central tendency ($\xi = 280.4122$) and variability ($\alpha = 268.5254$), with positive skewness ($k = 2.5539$), indicating right-skewed data and a higher likelihood of extreme flood magnitudes compared to GLO and GPA (Hamed & Rao, 2019).

Table 4. Performance measurement for candidate distributions

Distribution	GLO	GPA	PE3
MAE	29.23121014	23.89130881	25.91029352
MAPE	0.102461202	0.080253432	0.110348838
RMSE	53.52337859	40.20421989	39.08422233
RMSPE	0.33820537	0.100496606	0.133524466
R^2	0.950932826	0.97390899	0.975573279

Table 4 represent the performance of GLO, GPA and PE3 distributions using various metrics. The GPA distribution has the lowest MAE (23.8913) and MAPE (0.0803), indicating the most accurate predictions, while PE3 has the smallest RMSE (39.0842), showing better handling of large errors. GPA also leads in RMSPE (0.1005), outperforming GLO (0.3382) and PE3 (0.1335). However, PE3 has the highest R^2 (0.9756), suggesting it explains slightly more variance in flood data than GPA (0.9739) and GLO (0.9509). Overall, GPA excels in accuracy, while PE3 performs best in minimizing large errors and explaining variability.

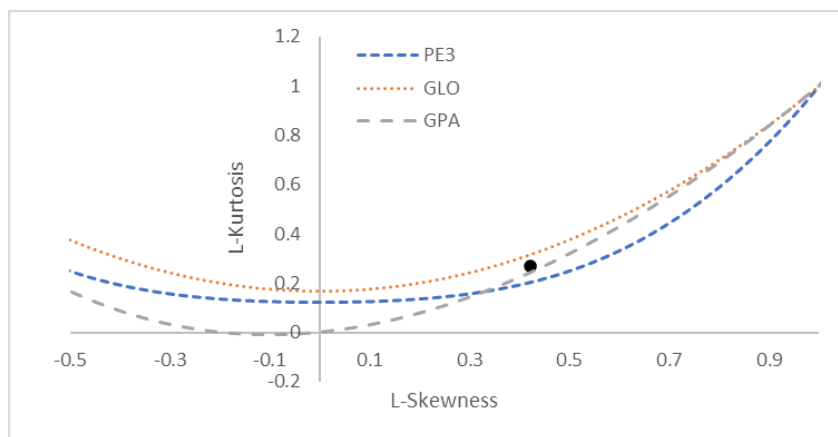


Fig. 3. L-moment ratio diagram

Fig. 3 displays the L-Moment Ratio Diagram (LMRD), comparing observed flood data which are plotted as a point with L-Skewness and L-Kurtosis = 0.2689 against theoretical curves of GLO, GPA and PE3 distributions. The observed point lies closest to the GPA curve, suggesting it may be the best fit, while

GLO and PE3 show greater deviation. To objectively evaluate fit, the Euclidean distance measures how closely each distribution's theoretical values match the observed data, providing a quantitative goodness-of-fit comparison between the three distributions.

Table 5. Euclidean distance for each distribution

Distribution	Euclidean Distance
GLO	0.04480754
GPA	0.028652209
PE3	0.0666223

Table 5 presents the Euclidean distance results comparing the fit of GLO, GPA, and PE3 distributions to the observed flood data. The GPA distribution demonstrates the strongest alignment with the smallest distance (0.0287), confirming it as the most accurate model for this dataset. The GLO distribution, while reasonable with a distance of 0.0448, falls short of GPA's precision. In contrast, the PE3 distribution shows the poorest fit (0.0666), indicating it is the least suitable option among the three. These results clearly prioritize GPA as the optimal choice for modelling flood events at this location, balancing theoretical suitability with empirical performance.

Table 6. Rank score for candidate distributions

Distribution	GLO	GPA	PE3
MAE	1	3	2
MAPE	2	3	1
RMSE	1	2	3
RMSPE	1	3	2
R^2	1	2	3
Euclidean Distance	2	3	1
Total Score	11	20	14

Table 6 ranks the candidate distributions (GLO, GPA, PE3) using a scoring system. For each test and measure, the best distribution gets 3 points, the worst gets 1, and the middle performer gets 2. The scores are totaled and the distribution that hold the maximum overall score is chosen as the most fit, showing the most reliable performance across all evaluations. The GPA distribution emerges as the best choice for flood frequency modeling in Johor's river systems. It excels in accuracy measures (MAE , $MAPE$, $RMSPE$, R^2) and has the smallest Euclidean Distance. With the highest total rank score in Table 6, GPA outperforms GLO and PE3, which showed weaker results. This consistent superiority across all tests confirms GPA as the most reliable option for flood risk assessment in the Johor River.

Table 7. Estimated flood discharge for the Johor River site

Return Periods (Years)	Probability (p)	Estimated Flood Discharges (m^3/s)		
		GLO	GPA	PE3
2	0.50	200.5983	189.757	182.494
10	0.90	532.7005	584.7902	614.3538
50	0.98	1107.315	1119.669	1102.259

Table 7 gives the return period estimate for each distribution. From the table above, the flood discharge estimates from the GPA distribution increase with longer return periods (2, 10, 50 and 100 years). For 2 years ($p = 0.50$), it's $189.757 m^3/s$, rising to $584.790 m^3/s$ (10 years, $p = 0.90$), $1119.669 m^3/s$ (50 years, $p = 0.98$), and reaching $1404.908 m^3/s$ for 100-year events ($p = 0.99$). This trend demonstrates GPA's effectiveness in modeling higher flood magnitudes for rare, extreme event.

4. CONCLUSION

The FFA is a widely recognized hydrologic engineering approach that has gained significant interest among researchers. This study focuses on identifying the most fitted probability distribution for the modelling of the annual peak flow data at a Kahang river site in Johor, Malaysia using the L-moment method for parameter estimation, L-moment ratio (LMR) analysis, and numerical performance criteria. The analysis relied on annual maximum flow data from the Kahang River's historical records. By employing multiple assessment tools, the study ensures the chosen distribution is reliable for predicting return periods. Results indicate that the GPA distribution best fits the annual peak flow data for this site, making it a strong candidate for both regional and local FFA applications in Johor's River systems for future research.

5. ACKNOWLEDGEMENTS/FUNDING

The authors acknowledged the financial support from the Ministry of Higher Education Malaysia through Fundamental Research Grant Scheme FRGS-EC/1/2024/STG06/UITM/02/12.

6. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

7. AUTHORS' CONTRIBUTIONS

Nur Hanida Othman: Conceptualization, Methodology, Formal analysis, Writing – original draft; **Basri Badyalina:** Supervision, Writing – review & editing, Validation, Project administration; **Ani Shabri:** Methodology, Validation, Resources, Writing – review & editing; **Muhammad Fadhil Marsani:** Data curation, Software, Visualization; **Fatin Farazh Ya'acob:** Investigation, Literature review, Data interpretation.

8. REFERENCES

- Adhami, M. (2024). Regional flood frequency analysis of Northern Iran. *Dokuz Eylül Üniversitesi Mühendislik Fakültesi Fen ve Mühendislik Dergisi*, 26(77), 272-280. <https://doi.org/10.21205/deufmd.2024267711>
- Ahmad, N. F., Fakhru'l-Razi, A., mat said, A., Azan, R., & Muhaimin, R. (2020). Community preparedness to flood disaster in Johor, Malaysia. *IOP Conference Series: Earth and Environmental Science*, 479(1), 012015. <https://doi.org/10.1088/1755-1315/479/1/012015>
- Ali, S. & Rahman, A. (2022). Development of a kriging-based regional flood frequency analysis technique for South-East Australia. *Natural Hazards*, 114(3), 2739-2765. <https://doi.org/10.1007/s11069-022-05488-4>
- Badyalina, B., Mokhtar, N. A., Mat Jan, N. A., Hassim, N. H. & Yusop, H. (2021). Flood frequency analysis using L-moment for Segamat river. *Matematika*, 37(2), 47-62. <https://doi.org/10.11113/matematika.v37.n2.1332>
- Badyalina, B., Shabri, A., & Marsani, M. F. (2021). Streamflow estimation at ungauged basin using modified group method of data handling. *Sains Malaysiana*, 50(9), 2765-2779. <https://doi.org/10.17576/jsm-2021-5009-22>

- Badyalina, B., Mokhtar, N. A., Ya'acob, F. F., Ramli, M. F., Majid, M., & Jan, N. A. M. (2022). Flood frequency analysis using L-moments at Labis in Bekok river station. *Applied Mathematical Sciences*, 16(11), 529-536. <https://doi.org/10.12988/ams.2022.917192>
- Che Ilias, I. S., Wan Zin, W. Z. & Jemain, A. A. (2021). Regional frequency analysis of extreme precipitation in Peninsular Malaysia, In *e-Proceedings of the 5th International Conference on Computing, Mathematics and Statistics (iCMS 2021)* (pp. 160-172).
- Fakhru'l-Razi, A., Diyana, F. A. N., Aini, M. S., Azan, R. A. & Muhaimin, R. W. M. (2020). Community preparedness to flood disaster in Johor, Malaysia. *IOP Conference Series: Earth and Environmental Science*, 479(1), 012015. <https://doi.org/10.1088/1755-1315/479/1/012015>
- Gringorten, I. I. (1963). A plotting rule for extreme probability paper. *Journal of Geophysical Research* (1896-1977), 68(3), 813-814. <https://doi.org/https://doi.org/10.1029/JZ068i003p00813>
- Hamed, K. & Rao, A. Ramachandro. (2019). *Flood frequency analysis*. CRC press. <https://doi.org/10.1201/9780429128813>
- Hassim, N. H., Badyalina, B., Mokhtar, N. A., Abd Jalal, M. Z. H., Zamani, N. D. B., Kerk, L. C. & Shaari, N. F. (2022). In situ flood frequency analysis used for water resource management in Kelantan River basin. *International Journal of Statistics and Probability*, 11(5). <https://doi.org/10.5539/ijsp.v11n5p8>
- Hosking, J. R. M. (1990). L-moments: analysis and estimation of distributions using linear combinations of order statistics. *Journal of the Royal Statistical Society: Series B: Statistical Methodology*, 52(1), 105-124. <https://doi.org/10.1111/j.2517-6161.1990.tb01775.x>
- Hosking, J. R. M. & Wallis, J. R. (1997). Regional frequency analysis. In J. R. M. Hosking & J. R. Wallis (Eds.), *Regional Frequency Analysis: An Approach Based on L-Moments* (pp. 1-13). Cambridge University Press. <https://doi.org/10.1017/CBO9780511529443.003>
- Jafry, N. A., Suhaila, J., Yusof, F., Nor, S. R. M. & Alias, N. E. (2023). Bivariate copula for flood frequency analysis in Johor River basin. *IOP Conference Series: Earth and Environmental Science*, 1167(1), 012018. <https://doi.org/10.1088/1755-1315/1167/1/012018>
- Jan, N. A. M., Shabri, A., Ismail, S., Badyalina, B., Abadan, S. S., & Yusof, N. (2016). Three-parameter lognormal distribution: Parametric estimation using L-moment and TL-moment approach. *Jurnal Teknologi (Sciences & Engineering)*, 78(6-11). <https://doi.org/10.11113/jt.v78.9202>
- Krislock, N., & Wolkowicz, H. (2012). Euclidean distance matrices and applications. In *Handbook on semidefinite, conic and polynomial optimization* (pp. 879-914). Springer. https://doi.org/10.1007/978-1-4614-0769-0_30
- Landwehr, J. M., Matalas, N. C. & Wallis, J. R. (1979). Probability weighted moments compared with some traditional techniques in estimating gumbel parameters and quantiles. *Water Resources Research*, 15(5), 1055-1064. <https://doi.org/10.1029/WR015i005p01055>
- Linsley, R. K. (1986). Flood estimates: How good are they?. *Water Resources Research*, 22(9S), 159S-164S. <https://doi.org/10.1029/WR022i09Sp0159S>
- Marsani, M. F., Shabri, A., Badyalina, B., Jan, N. A. M., & Kasihmuddin, M. S. M. (2022). Efficient market hypothesis for Malaysian extreme stock return: Peaks over a threshold method. *Matematika*, 141-155. <https://doi.org/10.11113/matematika.v38.n2.1396>
- Moughamian, M. S., McLaughlin, D. B. & Bras, R. L. (1987). Estimation of flood frequency: An evaluation of two derived distribution procedures. *Water Resources Research*, 23(7), 1309-1319. <https://doi.org/10.1029/WR023i007p01309>

- Romali, N. S. & Yusop, Z. (2017). Frequency analysis of annual maximum flood for Segamat river. In *MATEC Web of Conferences*, EDP Sciences, 103, 04003. <https://doi.org/10.1051/mateconf/201710304003>
- Stedinger, J. (1993). Frequency analysis of extreme events. *Handbook of Hydrology*, 18.
- Todd, D. K. (1957). Frequency Analysis of Streamflow Data. *Journal of the Hydraulics Division*, 83(1), 1166-1-1166-18. <https://doi.org/10.1061/JYCEAJ.0000060>
- Turhan, E. (2022). An investigation on the effect of outliers for flood frequency analysis: The case of the Eastern Mediterranean Basin, Turkey. *Sustainability*, 14(24), 16558. <https://doi.org/10.3390/su142416558>
- Valentini, M. H. K., Beskow, S., Beskow, T. L. C., de Mello, C. R., Cassalho, F. & da Silva, M. E. S. (2024). At-site flood frequency analysis in Brazil. *Natural Hazards*, 120(1), 601–618. <https://doi.org/10.1007/s11069-023-06231-3>
- Zamani, N. D., Badyalina, B., Jalal, A., Hakim, M. Z., Mohamad Khalid, R., Ya'acob, F. F. & Chang, K. L. (2024). Assessing flood risk using L-Moments: An analysis of the generalized logistic distribution and the generalized extreme value distribution at Sayong River Station. *Mathematical Sciences and Informatics Journal (MIJ)*, 5(2), 105–115.
- Zhang, J.-M., Chen, X.-H. & Ye, C.-Q. (2012). Calculation of negative-skewness hydrological series with Pearson-III frequency curve. *Journal of Hydraulic Engineering*, 43(11), 1296–1301. [https://doi.org/10.1061/\(ASCE\)HE.1943-5584.0000580](https://doi.org/10.1061/(ASCE)HE.1943-5584.0000580)



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

User Acceptance Testing of the 'Furwish' Location-Based Mobile Application: An Empirical Study on Pet Adoption and Surrender Services

Marnie Umirah Mazlan¹, Muhammad Nabil Fikri Jamaluddin^{2*}, Iman Hazwam
Abd Halim³, Alif Faisal Ibrahim⁴, Mohd Faris Mohd Fuzi⁵

^{1,2,3,4,5}Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Perlis Branch, Arau Campus,
02600 Arau, Perlis, Malaysia.

ARTICLE INFO

Article history:

Received 27 June 2025

Revised 3 August 2025

Accepted 6 August 2025

Published 1 September 2025

Keywords:

Pet Adoption

Pet Surrender

Animal Welfare

Location Aware

Mobile Application

DOI:

10.24191/jcrinn.v10i2.552

ABSTRACT

The increasing number of stray animals demands an effective way to connect pet owners who intend to surrender their pet with potential pet adopters. The manual approach to pet surrender and adoption often faced many challenges, such as limited outreach, communication, and options for the adopters. Unadopted pets may lead to an increasing number of stray animals and cause disturbance in residential areas. In this paper, we present a more detailed discussion on the development and evaluation of a location-aware mobile application named FurWish. The mobile application allows users to post information about their pet for listings, location-aware pet searches, setting up appointments for adoption, and notifications. The development adapted a three-phase methodology from Waterfall Model which includes requirement identification, design and development, and evaluation. Requirement identification involved reviews of literature and existing mobile apps. In the design phase the user interface design is drafted using diagrams and sketches, and the development phase involved implementation using Android Studio and Firebase with the integration of Google Location Services for displaying maps. In the evaluation phase, a user acceptance test (UAT) is conducted to test four key aspects, which are Attitude Towards Using (ATT), Behavioural Intention to Use (BI), Perceived Usefulness (PU) and Perceived Ease of Use (PEU). 33 respondents were given 12 questions after they had explored and used the functions and features of the mobile application. Results from UAT show FurWish has achieved high scores over four key aspects, with the highest score achieved under PU. This shows the user's acceptance of its effectiveness in supporting the pet adoption and surrender process while supporting animal welfare.

1. INTRODUCTION

Owning pet offers a few benefits on both pets and adopters which includes reducing stress level, companionship, and emotional support. A study shows that pet ownership can improve the levels of

^{2*} Corresponding author. E-mail address: nabilfikri@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.552>

human's physical activities (Martins et al., 2023). These positive effects encourage people adopting pets to be part of their life. Adopting pet usually comes with responsibility to make sure their nutritional needs are met and providing essential health care. Some pets require specialized diets and living environments to keep them physically and mentally healthy. It involved a continuous expense that the pet owners should consider on daily or monthly basis for maintaining their healthy pet. The costs incurred for maintaining pet may cause financial burden on certain pet owners. Therefore, owners who are unable to continue keep their pet should take responsibility on finding alternative such as surrendering the animal thru proper channels such as animal shelters and find adopters.

Pet owner who intends to surrender should not simply abandon and let it roams freely. The rising number of pets being abandoned contributes to the increasing number of stray animals. These abandoned animals may suffer from hunger, diseases and injuries. It also may cause disturbance in residential areas such as noise, spread of diseases, ransacked waste and safety issues among residents. The pet owner should be responsible to find the animal shelter or a new potential adopter to adopt the pet. Pet adoption is an important aspect of maintaining animal welfare and promotes responsible pet ownership within the communities. Traditional pet adoption methods faced challenges including limited outreach where many relies on physical visits and word of mouths. Problems in communication among pet owner and potential adopters also impairs the efforts in pet adoption. The other common challenge includes limited number of pets for selection to match their preferences and needs (Bhadane et al., 2023). These challenges may affect encouragement and number of participations in pet adoption.

To overcome the challenges, we developed a FurWish mobile application to support the process of adoption between current pet owners and potential pet adopters. The application developed implements the concept of location-aware which intended to improve the accessibility of available pets for adoption within the nearby locations. Users with the intention to surrender or adopt may view and post pet listings based on their geographical location. It also let users to set appointments and notification features to effectively facilitates the process of adoption. A brief introduction about FurWish mobile application was previously published as extended abstract that provide introductory aspects of the application and general overviews (Mazlan & Jamaluddin, 2025). In this paper we intend to present more detailed discussion and explanation about the design, development, features and results from the user's acceptance test.

The FurWish app is developed based on the adapted Waterfall Model, which considers suitable phases required for the goals of the mobile application development. Three phases involved are requirement identification, design and development and evaluation. Requirement identification includes identifying development feasibility of the pet adoption mobile application, review literatures and related works. Design phase considers drafting user interfaces and database design. Development phase involved development of the mobile application using Android Studio with integration of Firebase for data storage and Google Location Services for location-based features. The developed mobile application is evaluated on its usability and acceptance thru user acceptance test (UAT) in evaluation phase. Remaining sections of this paper presents detailed review of literatures, development process, features and findings from the evaluation.

2. LITERATURE REVIEW

2.1 Pet adoption

Stray animals refer to pet animals which are roaming freely on the streets without home and left behind by the owner. Their increase in population cannot be controlled, for example cats and dogs. These stray animals are considered as major public health problems (Abdulkarim et al., 2021). Stray animals, often facing difficulties surviving independently and at risk of various dangers such as hunger, harsh weather conditions, illnesses and traffic incidents. Various organizations and shelters are dedicated to rescue and

care of stray animals. Their main objective is to either reconnect the animals with previous owners or adopters.

Pet adoption plays an important role in reducing the rising number of stray animals and improving animals' welfare. It also may help in addressing overcrowding shelters and rescue organizations. There are also reports that small amount of unadopted pets in shelter would have to spend their entire life in shelter (Corsetti et al., 2025). Pet adoption is a process of a person to acquire a pet, whether to purchase or adopt the animal (Holland, 2019). Rather than purchasing pets from breeders or pet shops, adoption often involves rescuing animals that have been left, abandoned, surrendered or stray animals. The decision to adopt pet contributes to the efforts of animal welfare and supports the works of the shelters.

However, taking care of pet may incur some challenges among persons with physical ability limitations and financial problems (Meier & Maurer, 2022). Some pets require specific diets and living environments to keep them physically and mentally healthy. It involved a continuous expense that the pet owners should consider on a daily or monthly basis. The costs incurred for maintaining pets may cause financial burden on certain pet owners. Financial demands may arise especially during unexpected illness or accidents that requires treatments with veterinary (Chur-Hansen et al., 2008). These financial constraints faced by the pet owner often led them to surrender their pets. This also contributes to overcrowding shelters and stray animals. There is a need for a platform that can connect the pet owner who intends to surrender with the potential adopters.

2.2 Location-aware technology

Systems and applications which utilize global location data to enhance services, improve user experiences and activate functions based on a device or user's physical location are known as location-aware technologies (Schmidtke, 2020). These devices utilize various methods such as GPS, Wi-Fi positioning, cell tower triangulation, and Bluetooth for detecting and employing location data (Nord et al., 2002). These methods allow location-aware technologies to accurately pinpoint a device's position and provide relevant services. The functionality and user experience of mobile applications and services can greatly improve from this ability. Location awareness is an essential feature for the mobile applications nowadays (Akhigbe et al., 2022). Location-aware mobile applications collect, process and utilize geolocation data to provide context-aware information and services.

Various mobile applications across several industries successfully adopted location-based services to enhance user experience, services provided and decision making. In transportation industries, e-hailing platforms such as Grab and AirAsia MOVE (Md Isa et al., 2024) relies on location-based information in finding nearby passengers or drivers, determining routes and real-time location tracking (Yan et al., 2024). In food delivery applications like Foodpanda, GrabFood and ShopeeFood uses location-based information to connect nearby restaurants, delivery drivers and potentials customers (Zhao & Bacao, 2020). While in tourism and navigation, apps like Google Maps, Waze and FourSquare utilize geolocations information in providing location-based recommendations and suggestions on interesting nearby places, route guidance, traffic updates and trending places. These examples of location-aware integrations in mobile applications show that it is significant in improving functionality and efficiency by providing personalized and real-time information to the users.

2.3 Waterfall model

Waterfall Model is widely used and one of the oldest software development methodologies. It was introduced by Royce in 1970 and never proposed the term "Waterfall" in his paper (Royce, 1970). Diagram used to present the model actually seems like a waterfall and clearly shows each phase flows from one phase to other phases (Matković & Tumbas, 2010). The model introduced a systematic sequential phases where each phase must be completed before beginning the following phase (Khan et al., 2024). Due to its

clarity and ease of management, the model is widely accepted and used in software development projects with well-defined requirements.

In this project, we use an adapted version of Waterfall Model to suit the scope, timeline and nature of mobile application development. The word "adapted" is referring to the selected phases from Waterfall Model which are implemented in this project, while omitting other phases due to the limits of this project timeline and scope. The selected phases include requirement gathering, system design, implementation and testing. This approach may improve flexibility and following systematic development steps. It also accepts feedback from the development after user acceptance testing (UAT) is carried out. Therefore, the customized development methodology approach is aligned with the best practices and at the same time suits the project's constraints and scope.

2.4 Related works

Mobile applications named "Pet Adopter: Adopt Pets" allows user to have direct connections with individuals seeking to adopt pets and list their pet for adoption such as dogs, cats, rabbits and other animals (In Sequence Software LLC, 2022). Key features for this app include browsing pet listings, view user profiles, chat functionality and locating adoptable pet in specific areas. However, the unlimited chat requests are only available withing the subscription. "Pet Adoption" mobile application allows user to browse pet of different breeds with information such as owner's contact information or shelter (Friendsapp Listing, 2023). It also allows users to filter pet searches based on gender, size, age, breed, and location.

The PAWS Animal Welfare Society (PAWS) website provides detailed pet listings, information on adoption procedures, tools for donors and volunteers (PAWS Malaysia, 2025). There is also a section about recent news and updates about the community. With PetRehomer website, user can explore available pets and list pets for adoption that promotes direct pet adoption and rehoming (PetRehomer, 2025). Important features include comprehensive pet profiles, breed and regional search criteria. The platform focuses on decreasing number of animals in shelters while providing adoption and rehomer guidelines and advices. St. Hubert's Animal Welfare Center webpage is another pet shelters that pet rehoming services, short-term crisis care, comprehensive pet listings and resources on adoption procedures (St. Hubert's Animal Welfare Center, 2025).

The following Table 1 shows a comparison of features offered by different platforms. The user experience can be improved by using mobile apps with extra features like favourites lists and geolocation, such "Pet Adopter" and "Pets Adoption". While they offer greater accessibility, websites like "Paws," "Pet Rehomer," and "St Hubert's Animal Welfare Center" lack some of the more advanced features available in mobile apps. Each platform offers basic features like pet profiles and contact choices, which are crucial for communication and adoption process.

Table 1. Comparison of features among pet adoption and surrender applications

	Pet Adopters: adopt pets	Pets Adoption	Paws: Paws Animals Welfare Society	Pet Rehomer	St Hubert's Animal Welfare Center
Platform	Mobile App	Mobile App	Website	Website	Website
Search Filter	✓	✓	✓	✓	✓
Pet Profiles	✓	✓	✓	✓	✓
Contact Options	✓	✓	✓	✓	✓
Geolocation	✓	✓	✓		
Favourites List	✓	✓			
Community Program					✓

Source: Author's review of literature

<https://doi.org/10.24191/jcrinn.v10i2.552>

3. METHODOLOGY

The development of FurWish mobile application involved three phases adapted from Waterfall Model which are requirement identification, design and development and evaluation. It also consists of activities planned for each phase to achieve the development objectives, refer Fig. 1. Requirement identification is the first phase, which includes identifying development feasibility of the pet adoption mobile application, review literatures and related works. Second phase includes the user interface, database design and actual development of the mobile application. The developed mobile application is then evaluated on its usability and acceptance thru user acceptance test (UAT) in last phase.

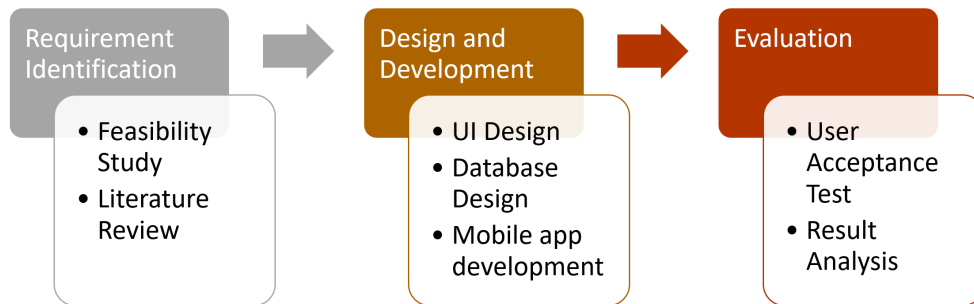


Fig. 1: Phases involving in development of FurWish mobile application

Source: Author's research methodology

3.1 Phase 1: Requirement identification

System requirements identification is the first step before starting project development. This phase is about figuring out what the project needs to achieve. A key part of this phase is the feasibility study, which reviews existing problems and solutions from previous works to find approaches that can be adopted in the mobile application. We collect information from various sources like websites, online libraries, journals, and expert opinions. This information includes details about the technologies and applications that can be used in the project. The feasibility study also defines the project's focus area, problems to be solved, objectives, scope and importance. By doing this, it ensures a clear understanding of what is needed to meet the project's goals. The subject matters reviewed include pet and animal adoptions, location-aware mobile applications, GPS, pet shelter and surrender.

3.2 Phase 2: Design and development

This phase is divided into two main activities that serve as crucial activities to achieve development goals of the FurWish mobile application. A few tools are utilized in preparing the design and development. Design activity includes creating drafts of user interfaces (UI) using Canva's wireframe tool based on functional requirements identified from requirement identification phase. The UI design also helps in visualizing layout and navigation flows of the mobile application. Database design involved preparing hierarchical data structure diagram using draw.io. It is used to model Firebase which is a NoSQL database model. The data structure includes five objects that include users, pets, adoptions, booking, notification and surrender. These objects are required to supports the applications user's authentication, database interactions and pet listings. Functional requirements identified are also modelled using use-case diagram to define interactions between users and the application. This also includes a flow chart to demonstrate the sequence of steps involved in surrendering and adopting pets.

The second activity in this phase is the development of the proposed mobile application. It involved actual implementation and development based on the design activities. The mobile application is developed using Android Studio IDE that focuses on implementing both application's front-end and back-end. The development also includes an integration of location-aware features by utilizing global positioning system (GPS) technology to track and utilize the real-time location of adopters. The location of the potential pet adopter is identified so that the application can display the list of nearby pet's owners who intended to surrender their pet. Using smartphone sensors, the application will query based on the current location of the potential adopters.

3.3 Phase 3: Evaluation

The last phase in methodology focused on evaluating the usability and overall acceptance of FurWish mobile application. User Acceptance Test (UAT) is conducted with a group of 33 selected respondents. Respondents include pet owners and non-pet owners with random genders and age ranges. UAT allows real users to interact with the application in practical scenarios, validating its usability and effectiveness in solving user problems. The respondents are listed with features to be tested that includes location-aware pet search, surrender pet via listings, setting appointments for adoption and notifications. Feedback from UAT is collected through questionnaire after respondents have used and tested the mobile application without any assistance and guidance by the developer. The questionnaire consists of 12 questions to evaluate four components of UAT which include Attitude Towards Using (ATT), Behavioral Intention to Use (BI), Perceived Usefulness (PU) and Perceived Ease of Use (PEU). The results obtained from the test are analyzed and discussed in this paper.

3.4 User interfaces and mobile application features

This section presents various user interfaces developed to suit various features offered by FurWish mobile application. Screenshot of the interfaces illustrate the overall user experience from registering as an authorized user towards adopting the pet. Upon launching the mobile application, user is displayed with "Login" screen, refer Fig. 2 (left) that asks for their registered email address and password. New users are provided with a link with "Signup Now" link that will direct users to "Sign Up" screen which can be depicted in Fig. 2 (centre). This screen allows new user to create an account by providing information including username, address, email and password. The creation of user's account is important to personalise user's experience and enable location-based pet searches. Upon successful registration, user will return back to "Login" screen.

Successful login will direct user to home screen of FurWish app as in Fig. 2 (right). The home screen features FurWish app's logo, access to notification and user's profile located at the top of the screen. There is an image banner about pet adoption to help users to better understand about the main purpose of the application. Right under the banner, there are four icons to access main features of the mobile application including "Adopt Pet", "Rehome Pet", "Appointment History" and "About". In this mobile application, users can act as adopter and owner, who intend to surrender their pet.

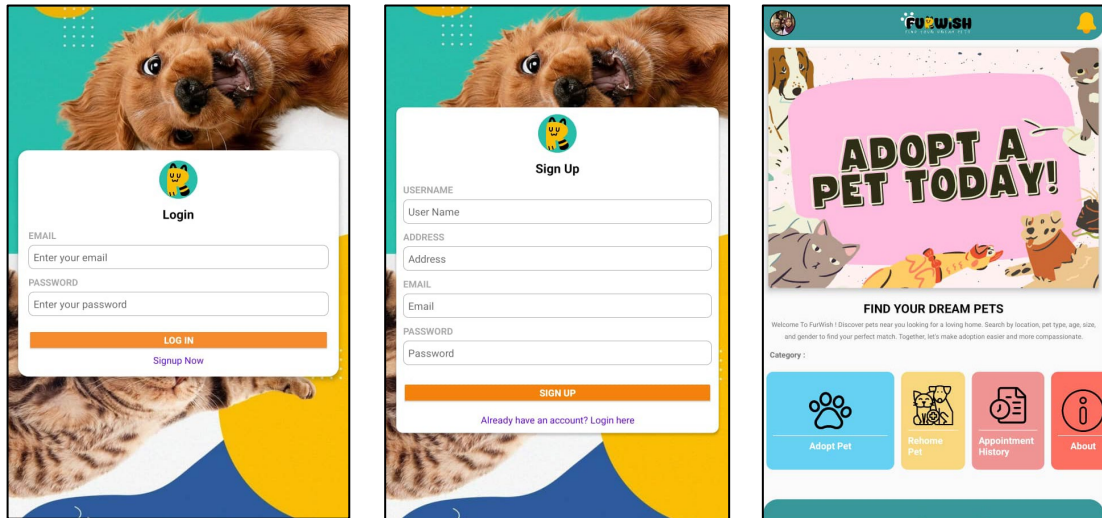


Fig. 2: (left) Sign up screen, (centre) login screen, (right) home screen

Source: FurWish mobile application screen shots

To begin with finding suitable pet for adoption, tapping on the button “Adopt Pet” will direct user to Adoption page as shown in Fig. 3 (left). This screen will first display list of all available pets to be adopted. User also has option to display list of available pets near to the user’s location by tapping on “Near You” button. This screen as in Fig. 3 (centre) features a map view with markers indicating the locations of the pets available for adoption, allowing user to click the markers to see the pet details. Below the map, detailed information about each pet is presented in a scrollable list including pet’s name, gender, age, photo, location and distance from user’s current position. Results displayed to user includes pets which are within five kilometers radius. Additionally, the screen includes a filter button at the top-right corner, which users can use to refine their search further based on user’s preference, such as type, gender or age, refer Fig. 3 (right).

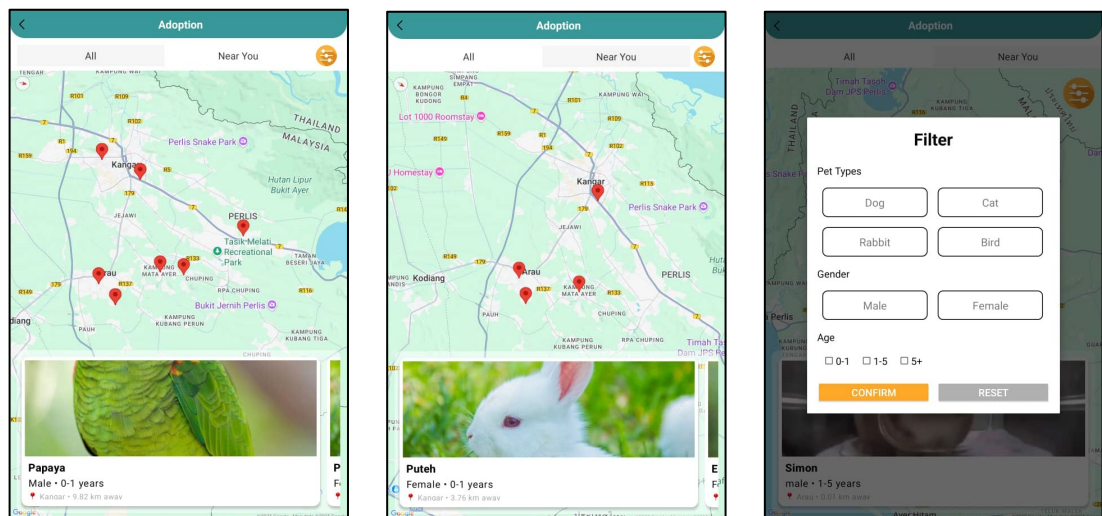


Fig. 3: (left) Listing all pets, (centre) listing of nearby pets, (right) pet filtering

Source: FurWish mobile application screen shots

Tapping on specific pet from the scrollable list will open the following screen as in Fig. 4 (left) where the mobile application will display the owners name and profile picture at the top section, followed by pet's image, name, gender, type, age, with its availability status. At the bottom, after the description, there is a map to display an approximate location of the pet and its owner's location together with an "ADOPT NOW" button which can be seen in Fig. 4 (centre). Tapping on the button will direct user to set an appointment for meeting with the pet owner. As can be seen in Fig. 4 (right), the mobile app requests user with appointment date and time with a message. Upon submission of the appointment date and time, the owner will be notified and asked for confirmation of the appointment.

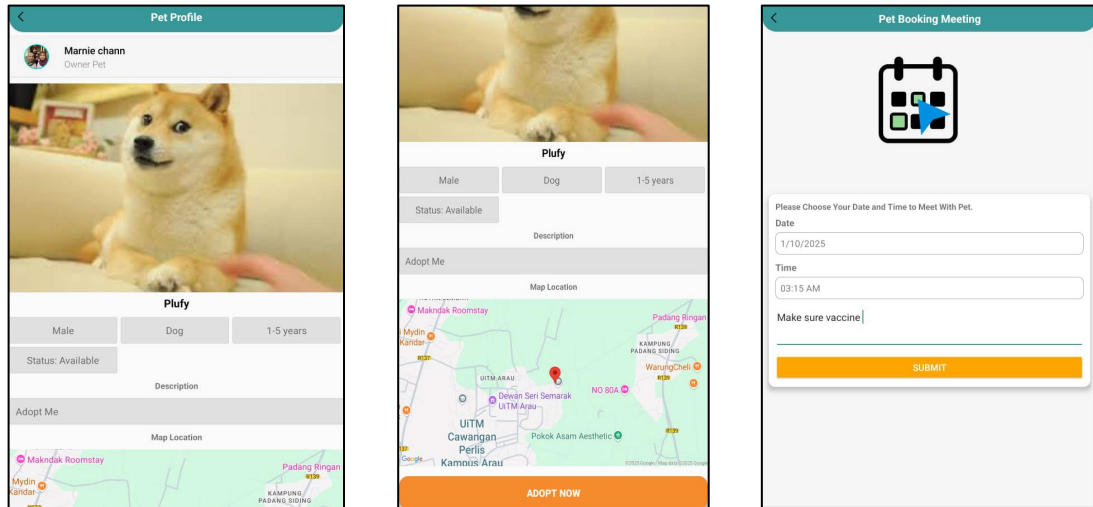


Fig. 4: (left) Pet information, (centre) pet location, (right) appointment submission form

Source: FurWish mobile application screen shots

The pet owner who intends to surrender or give away their pet can tap on "Rehome Pet" in the home screen of the FurWish mobile application. This will allow the owner to advertise their pet to potential adopter. The "Rehome Pet" button will direct user to the "Pet Details" screen as depicted in Fig. 5 (left). User is asked to upload an image of the pet to surrender, including its name, type such as Dog, Cat, Rabbit or Bird and information about gender and age. Additionally, users should write a brief description about the pet, highlighting any special traits or needs that potential adopter should know about. Once the form is completed and submitted using "POST" button, pet listing will be uploaded and visible to other users and potential adopters.

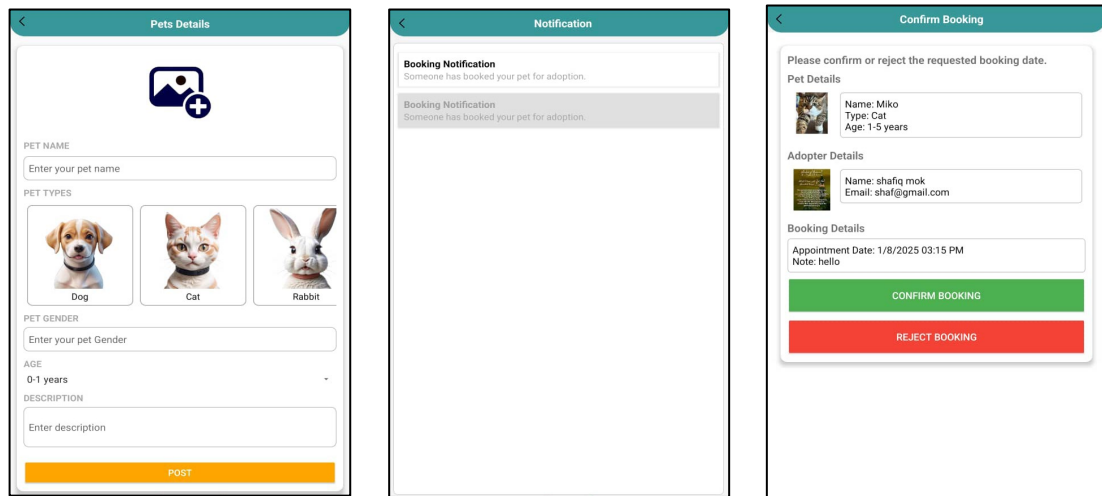


Fig. 5: (left) Rehome pet, (centre) app's notification, (right) booking confirmation

Source: FurWish mobile application screen shots

Notification screen of the FurWish mobile application can be accessed by tapping on the bell-shaped icon from the main screen. The notifications screen as can be seen on Fig. 5 (centre) provide users with important notifications about pet adoption activities, appointment requests, confirmation and interest from potential adopters. For example, it notifies when someone has booked their pet for adoption or express interest and ensures the user is promptly informed about any required actions to be taken. Unread notification appears in bold text while read notifications are displayed in regular and faded colour. Supposed user received a "Booking Notification", by tapping on the notification, the FurWish mobile app will display "Confirm Booking" screen (refer Fig. 5 (right)) asking the pet owner to accept or reject the adoption request for their pet. The screen also displays information about the pet, such as its name, type and age, together with adopter details which are name and email address. The booking details as requested by the potential adopter is also displayed which include appointment date, time and note from the adopter.

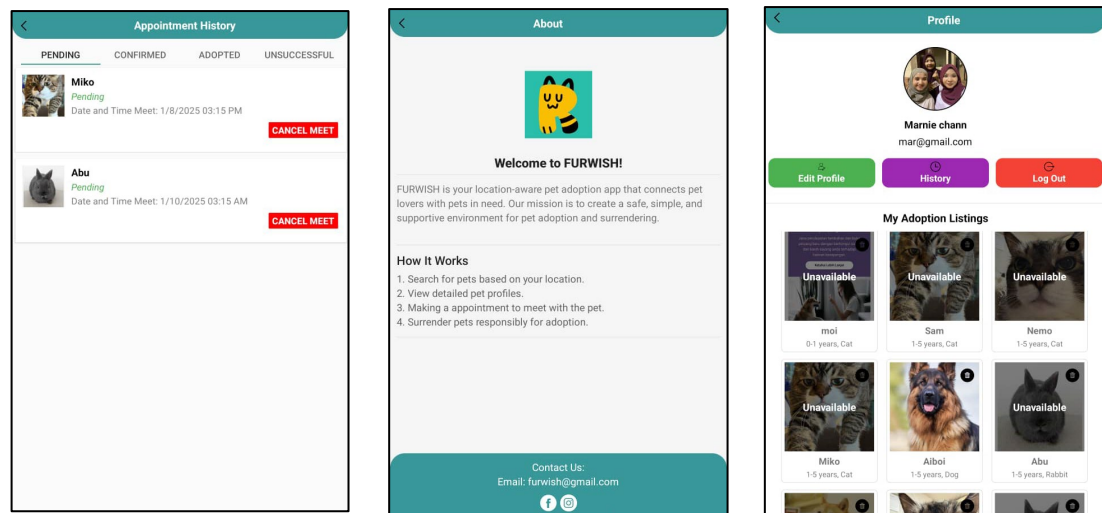


Fig. 6: (left) Appointment history, (centre) about, (right) history

<https://doi.org/10.24191/jcrinn.v10i2.552>

Source: FurWish mobile application screen shots.

At the home screen of FurWish mobile application, user can tap on “Appointment History” button to view list of adoption appointments as illustrated in Fig. 6 (left). The screen presents four tabs: “Pending”, “Confirmed”, “Adopted” and “Unsuccessful”. “Pending” tab lists appointment requests initiated by potential adopter which are waiting for the owner’s approval. The appointment request will remain in the tab until the owner accepts or rejects the request. It also will automatically cancel if no action is taken by the pet owner right after the meeting date and time. The potential adopter also has the ability to reject request by tapping on the “Cancel Meet” button. If pet owner accepts the appointment request, it will be moved to the “Confirmed” tab. After the meeting between pet owner and adopter, the adopter can tap on the “Adopt” button to confirm the adoption and the pet owner should confirm the adoption thru mobile application. The confirmed pet for adoption will then move to “Adopted” tab that lists all pets that have been successfully adopted. “Unsuccessful” tab lists cancelled or rejected appointments, in can be either owner rejects the appointment request or user cancels before the meeting date and time. Another icon accessible from main screen is “About” button which displays screen with welcoming message, introduction about the app and brief explanation on how it works, as shown in Fig. 6 (centre).

User’s profile screen can be accessed thru the main screen by tapping on the user’s icon at the top left corner. User’s profile screen as can be seen on Fig. 6 (right) displays user’s profile picture, name, email address at the top of the screen. Below the user’s information, there are three buttons to let user update user’s profile, view historical adoptions and log out of the application. Bottom section displays adoption listings made by user where the “Unavailable” marks indicate that it is in adoption process or have been adopted. Fig. 7 (left) displays an “Edit Profile” screen that allows users to update their username, address and email, while in Fig. 7 (right) shows two sections: “Pet Adoption History” and “Pet Surrender History”. First section lists all the adoption attempts made by the user, showing details like the pet’s name, type, age, gender, description, booking status, appointment date and notes. Second section shows pets that user has surrendered or offered for adoption through the app. It includes details like the pet’s name, type, age and the surrender status.

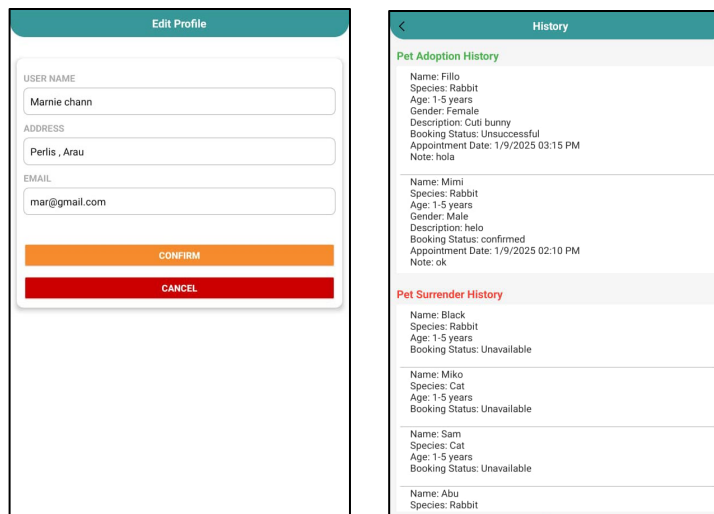


Fig. 7: (left) User’s profile, (right) edit profile

Source: FurWish mobile application screen shots

4. RESULTS AND DISCUSSIONS

This section presents results from the evaluation phase of FurWish development by focusing on the user's feedback and acceptance towards the mobile application. The results are based on collection of data from questionnaires through User Acceptance Test (UAT). The focus is on four key aspects of UAT which includes Attitude Towards Using (ATT), Behavioural Intention to Use (BI), Perceived Usefulness (PU), and Perceived Ease of Use (PEU) that make up 12 questions for evaluation. These questions are distributed over 33 randomly selected respondents. They are given time to explore and encouraged to use all features and functionality implemented in the mobile app. At the end of the session, respondents are requested to give their feedback via Google Form. The questionnaire utilized a five-point Likert scale to assess user perceptions and acceptance, where 1 indicated "Strongly Disagree" while 5 indicated "Strongly Agree".

Four key aspects that were evaluated in UAT is summarized in Table 2. PEU consists of four questions that related to how easy users could navigate and interact with the application. Two of the questions about easy of navigation and the process of surrendering pet had achieved 4.4 mean score. The remaining two questions about finding pet in the application and overall ease of use achieved slightly higher with 4.5 mean score. This resulted the overall mean score for PEU is 4.45, where in general most users found the mobile application is easy to use. Second item evaluated is PU that consisted of three questions which focused on how effective and helpful of the mobile application towards user in facilitating pet adoption process. Two questions about efficiency of adoption process and ability to discover potential adopters or pet owners achieved 4.5 average score, while one question related to the main goal of this mobile application which is location-aware feature achieved higher score with 4.6. The score 4.6 was also the highest score achieved among 12 questions distributed to respondents, this also leads to the overall best score for PU with 4.53 mean score.

Table 2. Summary of mean score from questionnaire.

Reaction of the Perceived Ease of Use (PEU)	Mean Score
I find the FurWish application easy to navigate.	4.4
I believe it is easy to find pets available for adoption on this application.	4.5
The process of surrendering a pet on the FurWish app is simple and straightforward.	4.4
Overall, I think FurWish is user-friendly.	4.5
Total Mean (PEU)	4.45
Reaction of Perceived Usefulness (PU)	Mean Score
FurWish helps make the pet adoption process more efficient.	4.5
Using the location-aware feature helps me find pets near my area.	4.6
The platform improves my ability to connect with potential adopters or pet owners.	4.5
Total Mean (PU)	4.53
Reaction of Attitude Towards Using (ATT)	Mean Score
I enjoy the idea of using FurWish for pet adoption or surrendering.	4.5
I have a positive attitude towards FurWish as a pet adoption platform.	4.5
Total Mean (ATT)	4.5

Reaction of Behavioural Intention (BI)	Mean Score
I would recommend FurWish to others who are interested in adopting or surrendering pets.	4.5
I plan to use FurWish again in the future if I need to adopt or surrender a pet.	4.5
Total Mean (BI)	4.5

The following is ATT, intended to assess user's feeling and attitude toward using the FurWish mobile application. The first questions evaluate the user's perception of adopting and surrendering pets thru a mobile application, while the second evaluate their overall positive attitude toward using the FurWish app. Both questions received an average score of 4.5, indicating strong user's acceptance and perception against the idea and features of the mobile application. The final aspect evaluated is BI, which measure the user's willingness to continue using the mobile application and recommending it to others. Both questions under this section also achieve an average score of 4.5. The result indicate that FurWish app developed is a practical tool for pet adoption and has intention to recommend to other users.

5. CONCLUSION

In conclusion, the FurWish mobile application is a successful location-aware mobile application that connects pet owners and potential adopters. It helps in preventing increasing number of stray animals while promoting animal welfare. The mobile application features location-aware pet search, surrender pet via listings, setting appointments for adoption and notifications. It also manages information for pet adoption and surrendering in facilitating user to achieve their expected goals. Results from User Acceptance Test (UAT) shows high level of acceptance over four different key aspects such as attitude towards using, behavioural intention to use, perceived usefulness and perceived ease of use. Highest total mean score obtained by PU indicates that user's acceptance on its effectiveness in supporting pet adoption and surrendering process. On the other hand, the lowest mean score obtained by PEU based on two questions obtained slightly lower value with 4.4 maybe due to the difficulty in interacting with the mobile application before they can get more familiar with it.

Future potential improvements may consider refining the flow of app's navigations and enhancing interface by addressing feedback from UAT. Other than that, the application also can integrate the connections with additional partnerships such as local pet shelters, rescue organizations and pet care providers. Additionally, real-time chat functionality can be added for enhanced communication between pet owners and adopters. Lastly, the integration with intelligence data processing such as machine learning and collaborative filtering algorithms for matching pet with user's preferences may enhances the overall user's experience.

6. ACKNOWLEDGEMENTS

The authors would like to acknowledge the support of Universiti Teknologi Mara (UiTM), Perlis Branch and Faculty of Computer and Mathematical Sciences, UiTM Perlis Branch for providing the facilities for this research.

7. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

8. AUTHORS' CONTRIBUTIONS

Marnie Umirah Mazlan: Conceptualisation, methodology, software development, formal analysis, investigation and writing-original draft; **Muhammad Nabil Fikri Jamaluddin:** Conceptualisation, methodology, formal analysis, validation, supervision, writing-review and editing; **Iman Hazwam Abd Halim:** Conceptualisation, methodology, formal analysis, writing-review and editing; **Alif Faisal Ibrahim:** Investigation, methodology and editing; **Mohd Faris Mohd Fuzi:** Methodology, writing-review and validation.

9. REFERENCES

- Abdulkarim, A., Goriman Khan, M. A. K. B., & Aklilu, E. (2021). Stray animal population control: Methods, public health concern, ethics, and animal welfare issues. *World's Veterinary Journal*, 11(3), 319–326. <https://doi.org/10.54203/scil.2021.wvj44>
- Akhigbe, B. I., Ososanya, O. J., & Olajubu, E. A. (2022). A mobile based pharmacy store location-aware system. *Journal of Information and Organizational Sciences*, 46(1), 213–236. <https://doi.org/10.31341/jios.46.1.12>
- Bhadane, D., Khirude, P., Chavan, O., & Lokare, A. (2023). Pet adoption system using web technology. *International Journal of Scientific Research in Engineering and Management (IJSREM)*, 07(05). <https://doi.org/10.55041/IJSREM22826>
- Chur-Hansen, A., Winefield, H., & Beckwith, M. (2008). Reasons given by elderly men and women for not owning a pet, and the implications for clinical practice and research. *Journal of Health Psychology*, 13(8), 988–995. <https://doi.org/10.1177/1359105308097961>
- Corsetti, S., Natoli, E., & Malandrucchio, L. (2025). The dark side of the moon: A good adoption rate conceals the unsolved ethical problem of never-adopted dogs. *Animals: an Open Access Journal from MDPI*, 15(5), 670. <https://doi.org/10.3390/ani15050670>
- Friendsapp Listing. (2023). *Pets Adoption - Adopt a Pet* [Mobile App]. App Store. <https://apps.apple.com/ph/app/pets-adoption-adopt-a-pet/id6450016352>
- Holland, K. E. (2019). Acquiring a Pet Dog: A review of factors affecting the decision-making of prospective dog owners. *Animals*, 9(4), 124. <https://doi.org/10.3390/ani9040124>
- In Sequence Software LLC. (2022). *Pet Adopter: Adopt pets* [Mobile App]. App Store. <https://apps.apple.com/my/app/pet-adopter-adopt-pets/id1268661502>
- Khan, R. A. H., Gopal, P. R., Yogesh, P. B., Nitin, P. V., Sahebrao, K. S., Sharad, P. N., & Santosh, P. V. (2024). The waterfall model-software engineering. *International Journal of Research Publication and Reviews*, 5(4), 7825–7830. <https://ijrpr.com/uploads/V5ISSUE4/IJRPR25788.pdf>
- Martins, C. F., Soares, J. P., Cortinhas, A., Silva, L., Cardoso, L., Pires, M. A., & Mota, M. P. (2023). Pet's influence on humans' daily physical activity and mental health: A meta-analysis. *Frontiers in Public Health*, 11, 1196199. <https://doi.org/10.3389/fpubh.2023.1196199>
- Matković, P., & Tumbas, P. (2010). A comparative overview of the evolution of software development models. *International Journal of Industrial Engineering and Management*, 1(4), 163–172. <https://doi.org/10.24867/IJIEEM-2010-4-019>
- Mazlan, M. U., & Jamaluddin, M. N. F. (2025). FurWish: Location-aware pet adoption mobile application. In *Proceedings of the Undergraduate Computer Sciences Colloquium 2025* (pp. 35–36).

- Md Isa, Y. B., Hilal Md Dahlan, N., Ezura, E., & Nisha Khirul Anuar, K. (2024). E-hailing services in Malaysia: A snapshot of legal issues and risks. *UCJC Business and Society Review (Formerly Known as Universia Business Review)*, 21(80). <https://doi.org/10.3232/UBR.2024.V21.N1.15>
- Meier, C., & Maurer, J. (2022). Buddy or burden? Patterns, perceptions, and experiences of pet ownership among older adults in Switzerland. *European Journal of Ageing*, 19(4), 1201–1212. <https://doi.org/10.1007/s10433-022-00696-0>
- Nord, J., Synnes, K., & Parnes, P. (2002). An architecture for location aware applications. In *Proceedings of the 35th Annual Hawaii International Conference on System Sciences* (pp. 3805–3810). IEEE. <https://doi.org/10.1109/HICSS.2002.994513>
- PAWS Malaysia. (2025). *Paws Animal Welfare Society (PAWS)* [Computer software]. <https://www.paws.org.my/>
- PetRehomer. (2025). *Home PetRehomer*. PetRehomer. <https://petrehomer.org/>
- Royce, D. W. W. (1970). Managing the development of large software systems. In *Proceedings of the 9th International Conference on Software Engineering* (pp. 328-338).
- Schmidtke, H. R. (2020). Location-aware systems or location-based services: A survey with applications to CoViD-19 contact tracking. *Journal of Reliable Intelligent Environments*, 6(4), 191–214. <https://doi.org/10.1007/s40860-020-00111-4>
- St. Hubert's Animal Welfare Center. (2025, March 5). *St. Hubert's Animal Welfare Center*. St. Hubert's Animal Welfare Center. <https://www.sthuberts.org>
- Yan, E. X., Md Jusoh, Z., Zainudin, N., & Osman, S. (2024). determinants of e-hailing service adoption among university students in Peninsular Malaysia. *International Journal of Academic Research in Business and Social Sciences*, 14(10), 1783-1803. <https://doi.org/10.6007/IJARBS/v14-i10/23281>
- Zhao, Y., & Bacao, F. (2020). What factors determining customer continuingly using food delivery apps during 2019 novel coronavirus pandemic period?. *International Journal of Hospitality Management*, 91, 102683. <https://doi.org/10.1016/j.ijhm.2020.102683>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Impact of Blackhole and Wormhole Attacks on DSDV Routing Protocol in VANET: Behavioural Analysis

Ahmad Yusri Dak^{1*}, Nuramarina Nasruddin², Nur Khairani Kamarudin³

^{1,2,3}*Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Perlis Branch, Arau Campus, 02600 Arau Perlis, Malaysia.*

ARTICLE INFO

Article history:

Received 16 March 2025

Revised 28 May 2025

Accepted 1 July 2025

Published 1 September 2025

Keywords:

VANET

Blackhole

Wormhole

DSDV

SUMO

DOI:

10.24191/jcrinn.v10i2.516

ABSTRACT

Vehicular Ad-Hoc Networks (VANETs) play a crucial role in Intelligent Transportation Systems (ITS) and the advancement of intelligent vehicles, enabling seamless and reliable communication between vehicles and infrastructure. This communication supports real-time applications such as collision avoidance, traffic control, and driver assistance. However, due to their dynamic topology and open-access medium, VANETs are highly vulnerable to security threats, particularly blackhole and wormhole attacks, which can disrupt data routing and severely degrade network performance. Despite the growing importance of VANETs, there is a limited number of studies that specifically examine the impact of blackhole and wormhole attacks on the Destination-Sequenced Distance-Vector (DSDV) routing protocol. This research addresses that gap by evaluating VANET performance under three conditions: no attack, a blackhole attack, and a wormhole attack. Key performance metrics including Packet Delivery Ratio (PDR), throughput, End-to-End Delay (EED), and Routing Overhead (RO) are analysed across various node densities using NS-2.35 and SUMO 1.18.0. Notably, the results show a 76.9% decline in throughput under the wormhole attack compared to the baseline scenario, highlighting the significant performance degradation caused by such threats. Overall, this study provides valuable quantitative insights into the vulnerabilities of VANETs and underscores the urgent need for more secure and resilient routing protocols to defend against these emerging attacks.

1. INTRODUCTION

In recent years, Vehicular Ad Hoc Networks (VANETs) have gained prominence in data transmission. They are increasingly used in Intelligent Transportation Systems (ITS) to support drivers, passengers, and services such as accident alerts and driver assistance systems (Mahmood et al., 2021). In addition, VANETs play a crucial role in safety applications by enabling communication between vehicles and roadside infrastructure. However, despite offering significant advantages for enhancing user safety, VANETs also

^{1*} Corresponding author. *E-mail address:* ahmadyusri@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.516>

pose various security challenges due to their open-access medium, high vehicle mobility, and dynamic topology changes. These characteristics result in a network topology that is highly dynamic and distinct, necessitating the use of efficient routing protocols. Reliable connectivity between vehicles and infrastructure is essential for effective communication, particularly under constantly changing network conditions, to ensure timely and accurate data transmission. In large-scale networks like VANETs, the Destination Sequenced Distance Vector (DSDV) protocol is effective as it offers loop-free routes using sequence numbers, employs proactive updates to prevent delays, and uses incremental updates to reduce overhead—advantages over other popular protocols such as AODV and DSR (El-Dalahmeh et al., 2024). In this study, the DSDV protocol periodically broadcasts its routing table to direct neighbours to maintain consistency in a rapidly changing topology (Sahoo & Tripathy, 2023). DSDV is based on the Bellman-Ford algorithm, where mobile nodes are required to broadcast their routing entries to nearby nodes (Alifo et al., 2023a).

Despite advancements in routing protocols for VANET security, inherent vulnerabilities persist, particularly concerning deception through false routing information disseminated by malicious nodes. These weaknesses can give rise to significant security threats, including blackhole and wormhole attacks, which compromise the reliability and integrity of VANET communications and hinder the implementation of real-time, dependable traffic management. A blackhole attack is a type of Denial of Service (DoS) in which a malicious node deceives other nodes by advertising the shortest path to a non-existent destination, thereby attracting and intercepting data transmissions. Once the data packets are routed through the attacker, they are simply dropped without notification, creating a "hole" in the communication pathway and preventing the data from reaching its intended recipient (Mahmood et al., 2021). Such disruptions can lead to delayed or lost safety-critical messages, such as collision alerts, ultimately degrading the performance and security of VANETs and increasing the risk of traffic congestion, accidents, and fatalities (Malik et al., 2022).

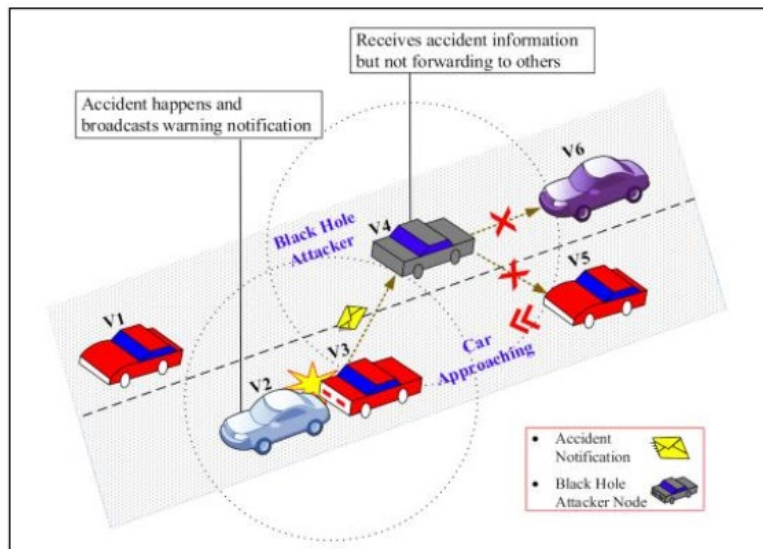


Fig. 1. Illustration of blackhole attack on VANET

Fig. 1 illustrates how a blackhole attack can affect communication in a VANET scenario. In this case, vehicle V3 collides with vehicle V2 and sends an alert message to vehicle V4, which happens to be a malicious node. Instead of forwarding the alert to nearby vehicles V5 and V6, vehicle V4 drops the message

entirely. This disruption could prevent other vehicles from receiving critical safety information in time, potentially leading to additional accidents or increased traffic congestion. In comparison, a wormhole attack is even more dangerous. Unlike the blackhole attack, which simply drops packets, the wormhole attack targets the routing layer and can bypass typical security measures such as cryptographic protections (Bhatti et al., 2024). It works by creating a tunnel between two colluding malicious nodes, allowing them to establish a shortcut in the network with a deceptively low hop count. This false route is then recorded in the routing tables, tricking legitimate nodes into using it (Asra, 2022). As a result, data can be intercepted, manipulated, or misrouted which leading to information theft, delays, or even complete breakdowns in communication (Bhatti et al., 2024). While these attacks don't block network channels directly, they severely disrupt the flow of information, affecting the core functionality and safety of VANET systems. Given these threats, this research seeks to simulate and analyse the performance of VANETs under blackhole and wormhole attacks using the DSDV routing protocol. The goal is to better understand how these attacks degrade network performance and to provide insights that can help automotive developers, researchers, and users design more secure and resilient communication systems for intelligent transportation.

2. RELATED WORK

In recent years, security has become a growing concern in both Vehicular Ad-Hoc Networks (VANETs) and Mobile Ad-Hoc Networks (MANETs), especially due to their vulnerability to various forms of cyberattacks. Among the most critical threats are blackhole and wormhole attacks, which can severely disrupt communication and degrade network performance. Several researchers (Al Rubaiei et al., 2022; Kaur & Kumar, 2018; Reshi et al., 2024) have explored these attack types and proposed different strategies to mitigate their effects. This section reviews recent studies that examine blackhole, wormhole, and greyhole attacks, with a focus on how well different routing protocols and detection methods perform under such conditions.

For example, Kumar et al. (2021) tackled the issue of blackhole attacks in VANETs by enhancing the AODV routing protocol. Their proposed solution involved improving the Route Request (RREQ) and Route Reply (RREP) processes and incorporating a cryptographic mechanism to verify node legitimacy. Using NS-2 simulations, they measured metrics such as Packet Loss Ratio (PLR), Packet Delivery Ratio (PDR), routing overhead, and End-to-End Delay (EED). Their findings showed that the enhanced protocol successfully detected and isolated malicious nodes, which led to improved delivery rates and overall network performance. In a similar line of study, Alifo et al. (2023b) evaluated how both AODV and Dynamic Source Routing (DSR) perform under blackhole attack conditions in MANETs. Their analysis revealed that AODV was more efficient in terms of energy consumption and delivery ratio, while DSR achieved better throughput. These results emphasize that the choice of routing protocol can significantly influence a network's resilience to attack, depending on the specific performance requirements. On the other hand, a study by Alifo et al. (2023a) focused on wormhole attacks and assessed their impact on DSDV, DSR, and TORA protocols in MANETs. The study found that DSDV and TORA were particularly vulnerable, with zero data delivery in some cases, while DSR offered relatively better protection, maintaining a modest PDR and higher throughput. This highlights the seriousness of wormhole attacks and the importance of developing robust mitigation strategies.

Another interesting approach was presented by Kaur et al. (2022), who addressed the less studied greyhole attack in VANETs. This type of attack involves selectively dropping or altering packets, making it harder to detect. The proposed detection mechanism used a three-phase method involving initialization, malicious node simulation, and system-wide broadcasts for detection. The approach led to significant improvements in both PDR and throughput, although it focused solely on greyhole attacks, leaving room for further research on more complex or combined threats. Lastly, Masruroh et al. (2022) studied blackhole and rushing attacks using the AOMDV routing protocol in a real-world traffic setting simulated with

OpenStreetMap, SUMO, and NS-2. Their findings showed that both attack types degraded Quality of Service (QoS), with blackhole attacks having a more severe impact. This included lower throughput and delivery rates, along with increased packet loss and delay. What sets this current study apart is its focus on the DSDV routing protocol in a VANET environment, which has been less commonly explored. Unlike prior work that often centers on AODV or DSR in MANET settings, this research uses a more VANET-specific approach by incorporating realistic urban mobility simulations through SUMO and NS-2. It also evaluates the performance of DSDV under both blackhole and wormhole attacks, across different node densities (10 to 50 nodes). Notably, the results reveal a dramatic 76.9% drop in throughput under wormhole attacks, an insight not often highlighted in previous research. By comparing attack scenarios with a baseline (no-attack) condition, this study offers a clearer picture of how DSDV holds up under real-world VANET challenges, providing valuable input for future protocol development and security enhancement. Table 1 summarises the key findings, objectives, and limitations of the related work reviewed in this study.

Table 1. Summary of Related Work

Author	Issue	Research Objectives	Limitation and Results
(Kumar et al., 2021)	VANET and traditional AODV routing protocols are susceptible to blackhole attacks.	To propose a secure AODV routing protocol to detect and mitigate blackhole attacks in VANET	Effectively detects and isolates malicious nodes, leading to a reduction in PLR.
(Alifo et al., 2023b)	Blackhole attacks in MANET can disrupt network services by dropping packets.	To evaluate the performance of AODV and DSR under blackhole attack	DSR shows higher throughput, while AODV has better PDR and residual energy.
(Alifo et al., 2023a)	Wormhole attacks can severely compromise the integrity and efficiency of data transmission in MANET.	To simulate and evaluate the performance of different routing protocols (DSDV, DSR, TORA) under wormhole attack.	- Zero data transmission and PDR in DSDV and TORA. - Higher throughput and PDR, but with a slight delay in DSR.
(Kaur et al., 2022)	Greyhole attacks randomly modify and discard packets making it difficult to detect in VANET.	To propose the detection and prevention mechanism of AODV routing protocol under greyhole attack in VANET	Effectively prevent the attack leading to a better result in PDR and throughput.
(Masruroh et al., 2022)	Blackhole and rushing attacks can disrupt VANET data communication.	To analyse and compare the effects of black hole and rushing attacks on VANET using AOMDV	Blackhole attack results in lower throughput and PDR, higher PLR, and increased EED compared to rushing attack.

3. METHODOLOGY

In this research, the NS-2 simulator is used to model a VANET environment and analyse its behavior under different scenarios. The simulations vary the number of nodes (10, 20, 30, 40, and 50) to represent different levels of traffic congestion within a network area of 1000×1000 square meters. Table 2 outlines the simulation parameters, providing a detailed overview of the setup, configurations, and experimental conditions. By using NS-2, the study aims to gain deeper insights into the performance characteristics of VANETs, particularly in the presence of blackhole and wormhole attacks. The selected parameters were chosen to reflect realistic urban conditions and are consistent with standard values commonly used in VANET research.

Table 2. Network Parameter

Parameter	Value
Simulator	NS-2.35
Mobility Model	SUMO Mobility (based on imported OSM map)
Routing Protocol	DSDV
Number of nodes	10, 20, 30, 40 and 50
Packet Size	1500 bytes
Simulation Time (s)	30, 60, 90, 120, and 150
Terrain Size (m)	1200 x 1400
Malicious Activity	Wormhole: 2 nodes Blackhole: 1 node
Performance Metric	EED, Throughput, PDR and Routing Overhead

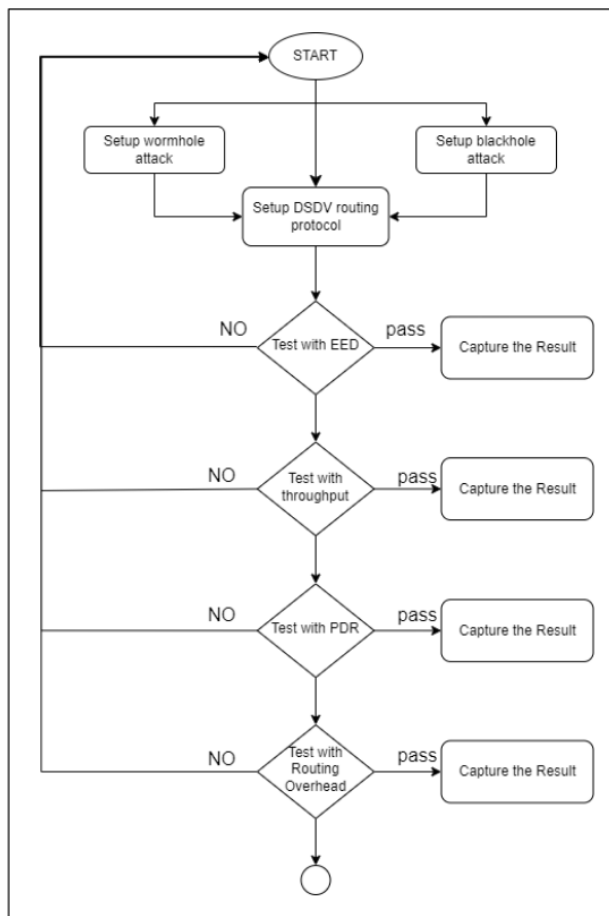


Fig. 2. Process flow for simulation phase

Fig. 2 shows the overall process used during the simulation phase of this research. To evaluate the network behaviour, each scenario was simulated using NS-2 for the networking components and SUMO to model realistic traffic flow and vehicle mobility. The research was carried out under three different

conditions: a baseline scenario with no attacks, a scenario involving a single blackhole node disrupting communication, and a third scenario where two wormhole nodes form a tunnel to manipulate data routing. The performance of each setup was assessed based on four main metrics: End-to-End Delay (EED), throughput, Packet Delivery Ratio (PDR), and routing overhead. Data collected from each simulation was then analysed to understand how these attacks impact overall VANET performance.

Fig. 3 displays the SUMO graphical user interface (GUI), which visualizes traffic flow and vehicle movements on a map imported from OpenStreetMap (OSM). This setup reflects realistic mobility patterns, helping to ensure the accuracy and relevance of the simulation results.

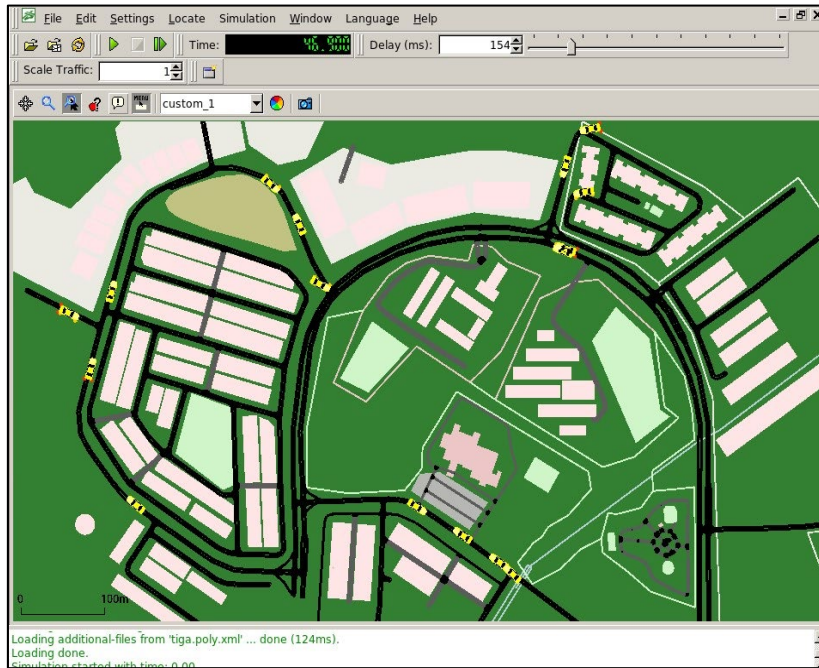


Fig. 3. SUMO-GUI

4. RESULT AND ANALYSIS

Results and analysis section is a critical part of this research, as it presents and interprets the findings from simulations aimed at evaluating the impact of blackhole and wormhole attacks on the DSDV routing protocol in a VANET environment. The analysis focuses on four key performance metrics: Packet Delivery Ratio (PDR), throughput, End-to-End Delay (EED), and Routing Overhead (RO). These metrics were chosen to offer a well-rounded view of the network's behaviour under both normal conditions and in the presence of malicious attacks. Results are presented using a comparative approach, contrasting the baseline performance of DSDV (without attacks) against its performance under blackhole and wormhole attack scenarios. This comparison helps to identify specific vulnerabilities within VANETs and assess the extent to which each type of attack degrades overall network performance.

4.1 Packet Delivery Ratio vs Number of Nodes

Fig. 4 illustrates how the Packet Delivery Ratio (PDR) changes as the number of nodes increases under three scenarios: without attack, with a blackhole attack, and with a wormhole attack. In the no-attack scenario, the PDR remains consistently high, ranging from approximately 98.7% to 99%, which reflects a stable and efficient network. Minor decreases in PDR at higher node counts can be attributed to increased congestion and packet collisions, but overall, the system performs reliably.

In contrast, both attack scenarios lead to noticeable reductions in PDR. The blackhole attack has a more significant impact, particularly as the network grows. At 30 nodes, for instance, the PDR drops to 95.56%, a decline of about 2.94% compared to the no-attack case (98.5%). This drop is due to the blackhole node attracting data by falsely advertising optimal routes and then discarding the packets, leading to direct data loss. Interestingly, the effect of the attack becomes less severe beyond 30 nodes, likely because the increased network density offers more alternate paths, reducing reliance on the malicious node. The wormhole attack also causes a decrease in PDR, but its impact is slightly less severe. At 30 nodes, the PDR falls to 95.55%, very close to the blackhole scenario. However, at 20 nodes, wormhole attacks result in a PDR of 97.35%, which is a smaller decline compared to the blackhole's 96.64%. This is because wormhole attacks typically involve tunneling packets between distant nodes, which disrupts routing paths but doesn't necessarily result in packet loss. As a result, some data still reaches its destination, maintaining a relatively higher PDR.

In summary, both attack types degrade network performance, but blackhole attacks are more damaging due to their tendency to drop packets entirely, while wormhole attacks primarily introduce routing inefficiencies that still allow some data to be delivered.

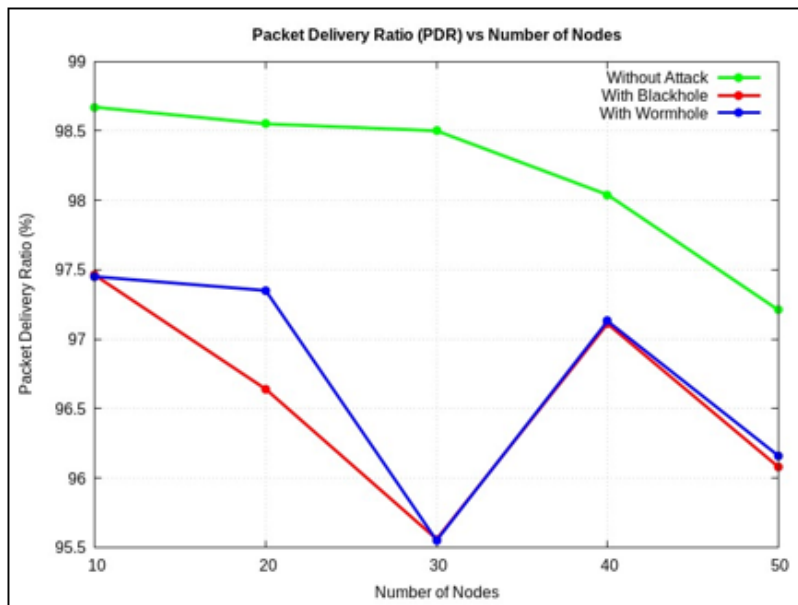


Fig. 4. Packet Delivery Ratio (PDR) vs Number of Nodes

4.2 Throughput vs Number of Nodes

Fig. 5 shows the variation in throughput as the number of nodes increases under three scenarios: without attack, with a blackhole attack, and with a wormhole attack. In the no-attack scenario, throughput

remains relatively high and stable, especially at lower node densities (10–20 nodes), where it reaches around 760 kbps. This performance reflects efficient data transmission and network operation under normal conditions, thanks to the proactive nature of the DSDV routing protocol. However, as node density increases, throughput declines to about 254.71 kbps due to greater congestion and competition for the shared wireless medium.

The blackhole attack has a dramatic impact, particularly at 20 nodes, where throughput drops sharply to 141.02 kbps, an 81.4% reduction compared to the baseline. This decline occurs because malicious nodes attract data traffic by advertising false routes, then drop the packets instead of forwarding them, leading to significant data loss. Although the impact is slightly less at other node levels, throughput consistently suffers under blackhole conditions. Wormhole attacks also degrade throughput, especially at 30 nodes, where it falls to 79.15 kbps, representing a 69% drop from the baseline. Unlike blackhole attacks, wormholes misroute packets through unauthorized shortcuts between colluding nodes, which causes routing inefficiencies and potential loops. Although some data still reaches its destination, the detoured paths consume more bandwidth and increase delays, thereby reducing overall throughput.

In summary, both types of attacks reduce network throughput, but blackhole attacks have a more severe effect due to direct packet drops, whereas wormhole attacks cause routing disruptions that moderately impact data delivery efficiency.

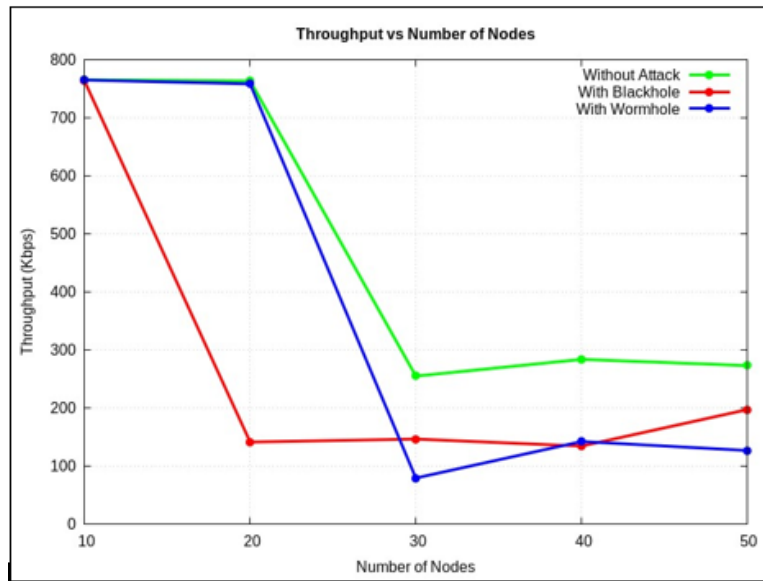


Fig. 5. Throughput vs Number of Nodes

4.3 End to End Delay (EED) vs Number of Nodes

Fig. 6 shows how End-to-End Delay (EED) changes as the number of nodes increases in three different scenarios: without any attack, with a blackhole attack, and with a wormhole attack. In the no-attack scenario, EED generally rises from 10 to 30 nodes, likely because more nodes in the network create congestion and delays as they compete to send data. Interestingly, the delay drops at 40 nodes, which might be due to a change in traffic flow or network structure during that simulation. However, at 50 nodes, the delay increases again, suggesting that higher congestion returns as the network becomes denser.

When a blackhole attack is present, the delay becomes noticeably higher—especially at 10 nodes, where it jumps from 132.193 ms to 181.262 ms, a 37.1% increase. This happens because data packets are drawn toward the malicious node and then dropped, forcing the network to try retransmitting the data, which takes more time. At 20 nodes, there's a slight and unexpected drop in EED compared to the no-attack case. This might be due to the blackhole node not actively participating in routing during that run, which could have reduced congestion by chance. While not typical, it shows how certain network conditions can sometimes reduce the visible impact of an attack. For the wormhole attack, the impact on EED is more mixed. At 10 and 20 nodes, there's a moderate increase in delay, as expected, because the attack disrupts normal routing. But at 30 nodes, the EED actually goes down. This could be because the wormhole creates a shortcut between distant nodes, reducing the number of hops and resulting in quicker data delivery in some cases.

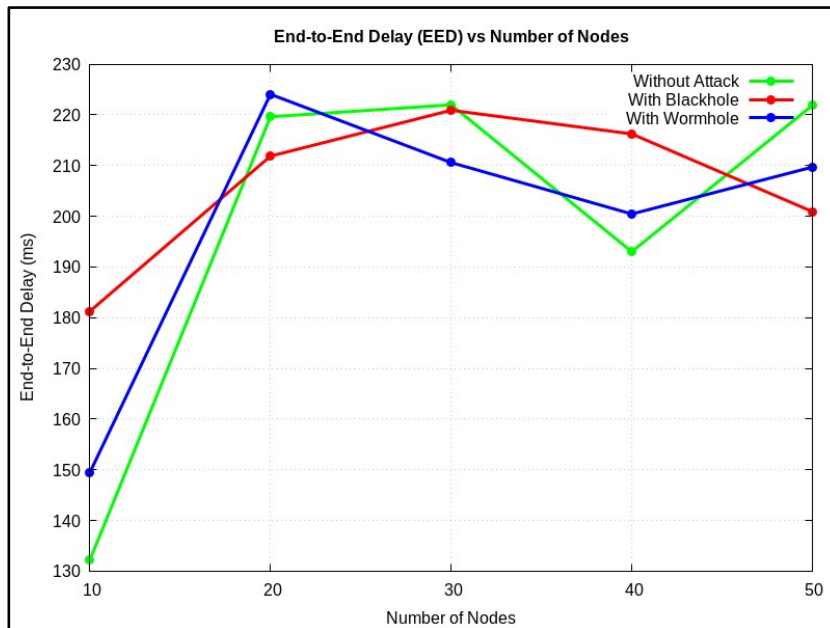


Fig. 6. End to End Delay (EED) vs Number of Nodes

4.4 Routing Overhead (RO) vs Number of Nodes

Fig. 7 illustrates how routing overhead changes with the number of nodes across different scenarios. In the no-attack scenario, the overhead increases gradually as the network grows. This is expected, as more nodes require more frequent routing updates and control messages to maintain connectivity and accurate routing tables. The DSDV protocol handles this growth efficiently, resulting in a steady but manageable rise in overhead, especially noticeable between 20 to 30 nodes and again from 40 to 50 nodes, reflecting the added complexity of a larger network.

When a blackhole attack is introduced, the routing overhead increases sharply. At 30 nodes, it spikes from 12.679 to 31.741, a jump of roughly 150 percent. This happens because the malicious node frequently advertises false routes, causing the network to constantly recalculate paths and send more control messages. Interestingly, at 50 nodes, the overhead actually drops below the no-attack level. This may be due to the network becoming dense enough to naturally route around the malicious node, reducing the disruption. The wormhole attack also raises routing overhead, but not as dramatically. At 30 nodes, it climbs to 22.317,

around a 76 percent increase, mainly due to routing inconsistencies and loops created by the attack's shortcuts. However, by the time the network reaches 50 nodes, the overhead from wormhole attacks becomes the highest of all three scenarios. This suggests that in very dense networks, the confusion caused by the wormhole grows more severe, leading to more control traffic as the protocol tries to resolve these routing issues.

In summary, both attack types disrupt normal routing and increase overhead. Blackhole attacks have a more immediate and dramatic impact due to frequent false route advertisements, while wormhole attacks cause more gradual but still significant increases, especially as the network becomes more complex.

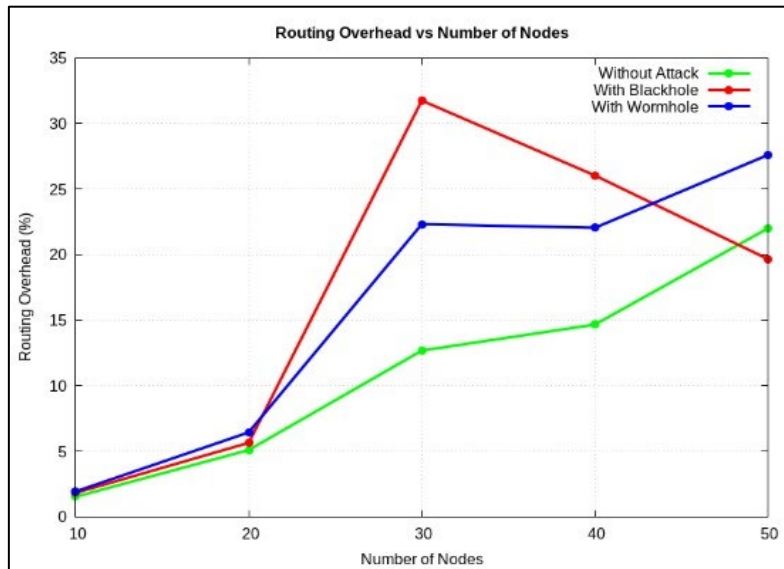


Fig. 7. End to End Delay (EED) vs Number of Nodes

5. CONCLUSION

This research provides meaningful insights into the performance and resilience of VANETs when subjected to blackhole and wormhole attacks, using the DSDV routing protocol as a baseline. The findings clearly show how these attacks can significantly impact key network metrics: namely Packet Delivery Ratio (PDR), throughput, End-to-End Delay (EED), and Routing Overhead (RO).

Under normal conditions, without any attacks, the network performs reliably. High PDR, stable throughput, and manageable delay and overhead highlight DSDV's effectiveness in maintaining communication in a dynamic VANET environment. However, when blackhole or wormhole attacks are introduced, this stability quickly breaks down. Blackhole attacks are particularly harmful, as malicious nodes drop packets outright, leading to sharp declines in PDR and throughput, increased delays, and a surge in routing overhead from constant path recalculations. Wormhole attacks, while slightly less destructive in terms of packet loss, still create serious disruptions by introducing false routes and increasing routing overhead, especially in denser networks. Interestingly, the results also show that blackhole attacks become somewhat less effective as the number of nodes increases, since the network can reroute data through alternative paths. In contrast, wormhole attacks tend to become more disruptive in larger networks due to the growing complexity of routing anomalies.

These findings highlight the urgent need for stronger security mechanisms in VANETs. By identifying how and where vulnerabilities occur, this research offers practical guidance for automakers, researchers, and developers aiming to build more secure and resilient vehicular networks. Future work could explore improvements to the DSDV protocol or assess alternative routing strategies better equipped to resist these types of attacks. Ultimately, this study reinforces the importance of balancing performance and security in the future of intelligent transportation systems.

6. ACKNOWLEDGEMENTS

The authors would like to acknowledge the support of Course Coordinator (Mrs Rafiza Ruslan), Faculty of Computer and Mathematical Sciences, Universiti Teknologi Mara (UiTM) Cawangan Perlis, Arau, Perlis, Malaysia for providing the facilities, guidance and moral support on this research.

7. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

8. AUTHORS' CONTRIBUTIONS

Ahmad Yusri Dak: Conceptualisation, methodology, formal analysis, visualization, investigation, data curation, writing-original draft, references; **Nuramarina Nasruddin:** Conceptualisation, funding acquisition, methodology, validation and writing – review & editing; **Nur Khairani Kamarudin:** validation and writing – review & editing.

9. REFERENCES

- Al Rubaie, M. H., Jassim, H. S., & Sharef, B. T. (2022). Performance analysis of black hole and worm hole attacks in MANETs. *International Journal of Communication Networks and Information Security (IJCNIS)*, 14(1). <https://doi.org/10.17762/ijcnis.v14i1.5078>
- Alifo, F., Doe, M., & Yakubu, M. A. (2023a). Wormhole attack vulnerability assessment of MANETs: Effects on routing protocols and network performance. *International Journal of Computer Applications*, 185(48), 24-29. <https://doi.org/10.5120/ijca2023923312>
- Alifo, F., Yakubu, M. A., Doe, M., & Asante, M. (2023b). Performance analysis of AODV and DSR routing protocols under blackhole attack using NS-2. *Computer Science & Information Technology*, 485-496. <https://doi.org/10.5121/csit.2023.131339>
- Asra, S. A. (2022). Security issues of vehicular ad hoc networks (VANET): A systematic review. *TIERs Information Technology Journal*, 3(1), 17-27. <https://doi.org/10.38043/tiers.v3i1.3520>
- Bhatti, D. S., Saleem, S., Imran, A., Kim, H. J., Kim, K., & Lee, K. (2024). Detection and isolation of wormhole nodes in wireless ad hoc networks based on post-wormhole actions. *Scientific Reports*, 14, 3428. <https://doi.org/10.1038/s41598-024-53938-9>
- El-Dalahmeh, M., El-Dalahmeh, A., & Adeel, U. (2024). Analysing the performance of AODV, OLSR, and DSDV routing protocols in VANET based on the ECIE method. *IET Networks*, 13(5-6), 377-394. <https://doi.org/10.1049/ntw2.12136>

- Kaur, G., Khurana, M., & Kaur, A. (2022). Gray hole attack detection and prevention system in vehicular adhoc network (VANET). In *3rd International Conference on Computing, Analytics and Networks (ICAN)* (pp. 1-6). IEEE. <https://doi.org/10.1109/ICAN56228.2022.10007192>
- Kaur, T., & Kumar, R. (2018). Mitigation of blackhole attacks and wormhole attacks in wireless sensor networks using AODV protocol. In *IEEE International Conference on Smart Energy Grid Engineering (SEGE)* (pp. 288–292). IEEE. <https://doi.org/10.1109/SEGE.2018.8499473>
- Kumar, A., Varadarajan, V., Kumar, A., Dadheech, P., Choudhary, S. S., Kumar, V. D. A., Panigrahi, B. K., & Veluvolu, K. C. (2021). Black hole attack detection in vehicular ad-hoc network using secure AODV routing algorithm. *Microprocessors and Microsystems*, 80, 103352. <https://doi.org/10.1016/j.micpro.2020.103352>
- Mahmood, J., Duan, Z., Yang, Y., Wang, Q., Nebhen, J., & Bhutta, M. N. M. (2021). Security in vehicular ad hoc networks: Challenges and countermeasures. *Security and Communication Networks*, 2021(1), 9997771. <https://doi.org/10.1155/2021/9997771>
- Malik, A., Khan, M. Z., Faisal, M., Khan, F., & Seo, J. (2022). An efficient dynamic solution for the detection and prevention of black hole attack in VANETs. *Sensors*, 22(5), 1897. <https://doi.org/10.3390/s22051897>
- Masruroh, S. U., Farghani, Y. S., KUSDARYONO, A., FIADÉ, A., PUTRI, R. A., & PRATIWI, L. A. (2022). Comparative analysis of testing black hole attack and rushing attack on VANET (Vehicular Ad-Hoc Network) with AOMDV routing protocol. In *International Conference on Engineering and Emerging Technologies (ICEET)* (pp. 1–6). IEEE. <https://doi.org/10.1109/ICEET56468.2022.10007272>
- Reshi, I. A., Sholla, S., & Najar, Z. A. (2024). Safeguarding IoT networks: Mitigating black hole attacks with an innovative defense algorithm. *Journal of Engineering Research*, 12(1), 133–139. <https://doi.org/10.1016/j.jer.2024.01.014>
- Sahoo, A., & Tripathy, A. K. (2023). On routing algorithms in the internet of vehicles: A survey. *Connection Science*, 35(1), 2272583. <https://doi.org/10.1080/09540091.2023.2272583>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

AI-Driven Forgery Detection in Offline Handwriting Signatures: Advances, Challenges, and the Role of Generative Adversarial Networks

Safura Adeela Sukiman^{1*}, Nor Azura Husin², Hazlina Hamdan³, Masrah Azrifah Azmi Murad⁴

¹*Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Johor Branch, Segamat Campus, 85000, Segamat, Malaysia.*

^{2,3,4}*Faculty of Computer Science and Information Technology, Universiti Putra Malaysia, 43400 Serdang, Malaysia.*

ARTICLE INFO

Article history:

Received 13 June 2025

Revised 19 July 2025

Accepted 22 July 2025

Published 1 September 2025

Keywords:

Offline Handwriting Signatures
Signature Forgery Detection
Generative Adversarial Networks
Siamese Networks
Autoencoders
Deep Learning

DOI:

10.24191/jcrinn.v10i2.532

ABSTRACT

Handwriting-based authentication continues to be a critical element in forensic analysis, particularly in the context of document fraud and signature forgery. Although deep learning (DL) techniques have shown promising results, there are still obstacles associated with the availability of limited datasets, the generalization of models, and their robustness. This review conducts a systematic examination of recent developments in DL methods for signature forgery detection. It employs the PRISMA protocol and retrieves literature from four well-established databases: Scopus, ACM Digital Library, Web of Science, and IEEE Xplore. Following a rigorous screening procedure, a total of 15 primary studies published between 2019 and 2025 were selected from an initial 115 records that were filtered by Computer Science subject area, English language, and original research articles. Five publicly accessible datasets: CEDAR, BHSig260, ICDAR 2011 SigComp, Kaggle signature verification dataset by RobinReni, and Kaggle handwritten signatures by Divyansh Rai were identified and analysed. The review indicates that Siamese networks dominate the DL architecture for signature forgery detection tasks, while alternative methods either employed fine-tuned pre-trained models (i.e., VGG16) or a hybrid of autoencoders and Convolutional Neural Networks (CNNs). An accuracy of 100% has been achieved through utilization of Siamese network leveraging the CEDAR dataset. This result is reasonable since CEDAR has the advantages of clean and balanced dataset. In response to the persisting limitations, this review emphasizes Generative Adversarial Networks (GANs) as the powerful data augmentation technique and a potential solution to enrich training datasets, simulate diverse forgery patterns, and enhance the robustness of models. Finally, a generative-aware conceptual framework is proposed at the end of the review to inform future research on the development of offline handwriting signature forgery detection system that is more resilient and forensic-ready.

^{1*} Corresponding author. *E-mail address:* safur185@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.532>

1. INTRODUCTION

Handwriting continues to play an important role in personal authentication, particularly in forensic and legal contexts where signature verification is central to identity validation and document security. Despite the proliferation of digital alternatives, handwritten signatures remain widely utilized in sectors such as banking, contracts, education, and healthcare, making them susceptible to forgery and fraud.

Offline handwriting signature forgeries can be broadly classified into three categories: random (or blind) forgeries, simple (or unskilled) forgeries, and skilled (or simulated) forgeries. Random forgeries occur when the forger has no access to the genuine signature and attempts to replicate it without any reference, resulting in signatures with little to no similarity to the original. Meanwhile, simple forgeries occur when the forger knows the person's name but does not have access to the actual signature, resulting in attempts that may mimic the general style but lack precision. Skilled forgeries, on the other hand, are produced by those who have access to actual signature samples and have practiced copying them, making them the most difficult to detect since they are so identical to original signatures (Hafemann et al., 2017).

The conventional methods of signature forgery detection, which frequently depend on human examiners or handcrafted features, are constrained in their scalability, objectivity, and accuracy. The automation of signature forgery detection has gained more prominence as a result of the emergence of artificial intelligence (AI) and DL. Sequential CNNs (Balaji et al., 2024), Siamese networks (Joe Harris & Anitha, 2023) and Autoencoder-based (Swamy et al., 2024) models have all shown significant potential in the detection of skilled and simulated forgeries. These models provide adaptability across writer-independent scenarios in addition to enhanced performance. Nevertheless, substantial challenges continue to exist, despite these advancements. The generalizability of the model is still constrained by dataset limitations, varying handwriting styles, and the complexity of skilled forgeries. Furthermore, most models trained on benchmark datasets often struggle to maintain robustness under real-world conditions that include noise, document degradation, and cross-domain variations (Engin et al., 2020).

Addressing these dataset limitations and generalizability issues, GANs have emerged as a powerful data augmentation component within the training pipeline. In contrast to discriminative models such as CNNs, which generally utilize conventional data augmentation methods like rotation, flipping, or scaling (Swamy et al., 2024), GANs is capable to generate entirely new and realistic samples that closely resemble the training distribution (Goodfellow et al., 2017; Wang & Jia, 2019). In the context of offline signature forgery detection, GANs can simulate high-fidelity synthetic forgeries, particularly skilled forgeries, which traditional augmentation methods often struggle to realistically re-create. This not only expands the dataset but also introduces controlled variability, allowing models to better generalize and resist overfitting. Consequently, GANs represent a significant paradigm shift in the approach to data scarcity, variability, and robustness in the context of AI-driven signature forgery detection. This review emphasizes GANs not merely as a supporting technique but as a transformative tool for enhancing the realism, diversity, and forensic-readiness of training datasets.

While several past reviews have examined handwriting and signature verification systems, many have predominantly concentrated on machine learning techniques (Soelistio et al., 2021) without incorporating the most recent advancements in deep learning and generative adversarial networks. Hafemann et al. (2017) conducted a survey of offline signature verification approaches but omitted current architectural innovations, like Transformer-based models and Siamese Networks. In contrast, this review incorporates recent peer-reviewed works (2019–2025) employing current deep learning architectures with multiple benchmark datasets. More importantly, it uniquely explores the emerging role of generative adversarial networks in addressing persistent challenges such as data scarcity, complexity, and model robustness, which remains underexplored in existing literature. The following shows the objectives of this review:

- a) To synthesize recent deep learning models in AI-driven offline handwriting signature forgery detection.
- b) To analyse challenges especially the ones related to offline handwriting signature datasets.
- c) To explore the emerging role of generative adversarial networks in addressing offline handwriting signature datasets challenges.
- d) To propose a generative-aware conceptual framework that is recent, more resilient and forensic-ready for future development of offline handwriting signature forgery detection system.

2. METHODOLOGY

The PRISMA (Preferred Reporting Items for Systematic Reviews and Meta-Analyses) protocol, a widely recognized framework for improving the rigor and reproducibility of systematic reviews in scientific domains, is invoked in this review to ensure a structured and transparent approach (Page et al., 2021). This protocol was modified to accommodate the scope of this review, which is focused on the application of deep learning in handwriting forensics, with a particular emphasis on the detection of offline signature forgery and identity-related manipulation. The PRISMA-based strategy facilitated a rigorous process of identification, screening, eligibility evaluation, and inclusion of relevant articles, thereby ensuring consistency and objectivity throughout the review process.

The review process consisted of four important stages: (a) the identification of relevant records from selected databases, (b) the removal of duplicates, (c) the assessment of titles and abstracts with respect to inclusion and exclusion criteria, and (d) the full-text evaluation of eligible studies.

2.1 Database selection and search strategy

Four established and high-impact databases were utilized for literature retrieval: Scopus, ACM Digital Library, IEEE Xplore, and Web of Science. The databases were chosen for their comprehensive coverage of peer-reviewed publications in the fields of computer science, artificial intelligence, and digital forensics. The following search queries were composed using meticulously structured Boolean expressions that combined domain-specific terms:

("handwriting forgery" OR "signature forgery" OR "document fraud") AND ("deep learning" OR "neural networks" OR "generative artificial intelligence" OR "generative adversarial networks")

The search was confined to articles published between 2019 and 2025, written in English, and classified under the Computer Science subject area. This period was chosen to indicate the time at which generative adversarial networks started to be practically utilized in handwriting forensics tasks including synthetic handwriting generation, anomaly detection, and counterfeit detection. Restricting the review to this time ensures that the included studies fit the present methodological setting and are pertinent to the state-of-the-art advancements in the domain. Reviews, editorials, and non-peer-reviewed works were excluded, and only original research articles and conference papers were considered.

2.2 Study identification and screening process

A preliminary collection of 115 records was obtained: Scopus (n = 47), ACM (n = 18), IEEE Xplore (n = 37), and Web of Science (n = 13). Following the elimination of 59 duplicate entries, 56 distinct records underwent a multi-phase screening process that included title screening, abstract screening, and full-text assessment (before any screening process n = 56).

The relevance of the articles was initially determined by looking at their titles. A total of eleven records were disregarded because they were deemed irrelevant. These records included studies that concentrated on hardware implementation or emotion detection rather than forgery analysis. After filtering the titles, the total number of articles is 45 (title screened n=45).

Following that, the relevance of articles was screened by reading their abstracts. Nine studies were disqualified because they presented models that were simply conceptual in nature, frameworks that were end-to-end but did not have empirical validation, or content that was not in English. The total number of articles after the screening of abstracts (abstract screened $n = 36$).

Finally, the inclusion criteria were met by 15 articles after the full texts were evaluated. The remaining 21 articles were excluded due to the following reasons: the use of privately collected or real-time data, the absence of performance evaluation, or the presentation of conceptual frameworks that were not reproducible. The total number of articles after full-text screening is 15 (full-text screened $n = 15$). Fig. 1. depicted the flow diagram of PRISMA protocol for both identification and screening process of this review.

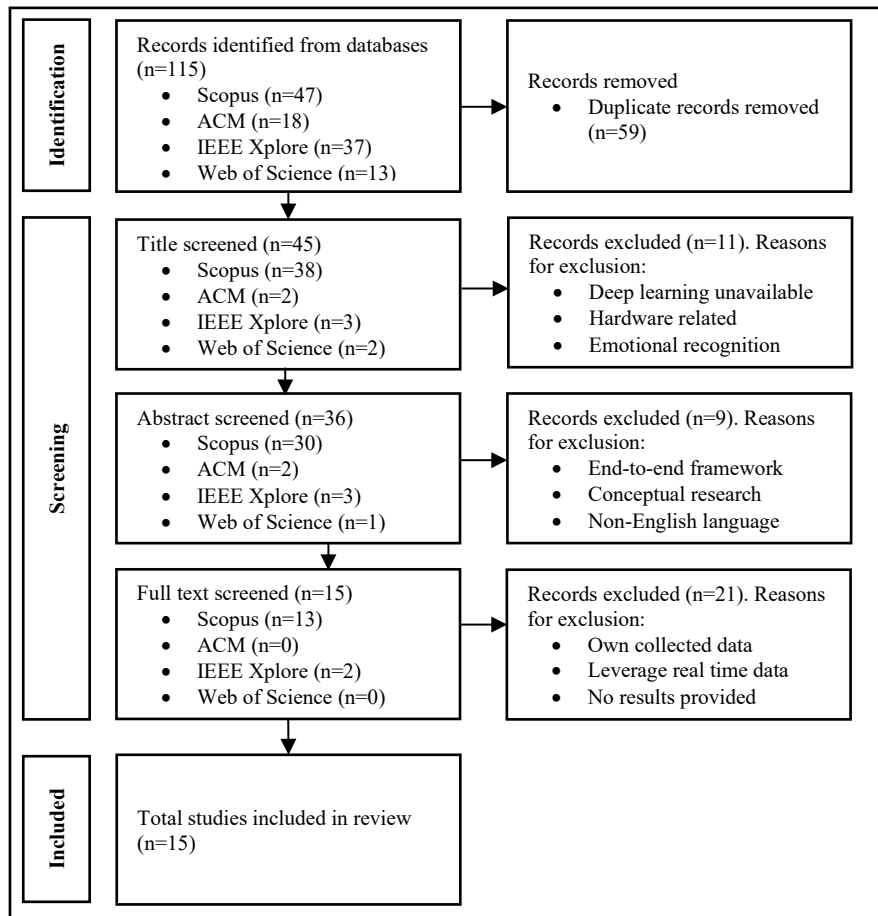


Fig. 1. PRISMA protocol flow diagram for filtering 15 relevant articles

2.3 Inclusion and exclusion criteria

This study employed inclusion and exclusion criteria to maintain the relevance and quality of the literature examined during the article selection process. Studies were selected if published in peer-reviewed journals or conference proceedings and focused on offline handwriting signature forgery detection. Moreover, qualifying studies must have utilized artificial intelligence (AI) or DL methodologies, encompassing, but not restricted to CNNs, Transformer-based architecture, hybrid models, or generative models. A fundamental requirement was the reporting of evaluation metrics pertinent to forgery detection,

including accuracy, precision, recall, and F1-score. Only research published in English were included to ensure consistency and interpretative accessibility.

On the other hand, studies were excluded if they were based on non-peer-reviewed sources, such as blog entries, technical reports, or preprints that did not undergo rigorous peer review. Articles that were exclusively theoretical, which lacking experimental validation or implementation results were also excluded. In addition, this review did not include research that exclusively focused on handwriting recognition (i.e., character or word transcription) without a direct emphasis on forgery or forensic detection applications.

Consequently, the following key attributes were extracted from each study: the specific forgery scenario addressed, the year of publication, the dataset(s) used, the learning model architecture, the use of generative techniques, performance results, and limitations noted by authors. This data was collected to facilitate comparative synthesis, trend analysis, and gap identification in the subsequent sections.

3. RESULTS

The results section is organized to initially provide an overview of the datasets, followed by the performance comparative analysis of all 15 deep learning models from the PRISMA protocol in offline handwriting signature forgery detection.

3.1 Dataset overview

This review incorporates five (5) publicly available offline signature datasets, which are extensively used in offline handwriting signature forgery detection studies. These datasets contain a balanced mix of genuine and forged samples from a wide range of signers and serve as important benchmarks for training and evaluating deep learning models. The datasets chosen include CEDAR (Srihari & Leedham, 2001), BHSig260 (Kathuria, 2022), ICDAR 2011 SigComp (Liwicki et al., 2011), Kaggle signature verification dataset by RobinReni (Reni, 2019), and Kaggle handwritten signatures by Divyansh Rai (Rai, 2020).

The CEDAR dataset (Srihari & Leedham, 2001) is widely regarded as a fundamental benchmark in offline signature verification. The dataset comprises 2,640 grayscale signature images from 55 signers, with each signer providing 24 genuine and 24 forged samples. A primary advantage is its balanced architecture and regulated acquisition settings, rendering it optimal for training and assessing models in a noise-free setting. Nonetheless, a significant disadvantage of the dataset is its comparatively small participant pool, which constrains its capacity to generalize across varied writing styles. The dataset consists solely of English-language signatures, hence constraining its relevance in multilingual or cross-script research.

Meanwhile, BHSig260 dataset (Kathuria, 2022) is a large-scale signature dataset with 260 signers: 100 in Bengali and 160 in Hindi. It comprises 24 authentic and 30 counterfeit signatures per user, amounting to almost 14,000 samples. Its strength is its wide range of languages and big sample size, which make it a useful benchmark for evaluating signature verification systems in scripts other than Latin. However, its constraints include script specificity, making it less relevant for Latin-based handwriting models as well as diversity in the quality and style of forgeries, which can generate errors during training or evaluation.

On the other hand, the ICDAR 2011 SigComp dataset (Liwicki et al., 2011) was initially developed for a competition. The signature images of 106 individuals are included in this dataset, with 64 of them being used for training and 42 for testing. The test set contains a minimum of 12 genuine and 12 forged samples per signer, while the training set contains a minimum of 24 genuine and 8 forged signatures per user. The dataset's primary advantage is its design for benchmarking, which allows for a rigorous writer-independent evaluation due to appropriate separation between training and testing users. Nevertheless, its constraints include the fact that it is relatively older and has some image quality constraints, as well as limited cultural or linguistic diversity, as it exclusively studies Dutch signatures.

Reni (2019) created the Kaggle signature verification dataset, which is a preprocessed version of Dutch signature samples from the ICDAR 2011 SigComp. It contains 2,149 labeled images of genuine and forged signatures. Its key advantages are accessibility and usability. This dataset is curated in a format that simplifies loading and training, making it ideal for rapid prototyping and model development. On the downside, it does not adhere to a rigorous train/test split and has a moderate dataset size, which may limit its appropriateness for models that require larger-scale evaluations.

Lastly, the Kaggle handwritten signatures of Divyansh Rai (Rai, 2020). This small collection has 300 signature images from 30 signers, with each providing five genuine and five forged samples. Its tiny size makes it ideal for educational applications, rapid model prototyping, and resource-constrained scenarios. However, its scale has severe limitations in terms of unpredictability, generalizability, and real-world use. It might not be appropriate for training deep models or conducting large-scale studies.

Collectively, these datasets encompass a wide variety of forgery types, writing scripts, and signer diversity, rendering them appropriate for assessing the robustness and efficacy of deep learning models in the detection of offline signature forgeries. Their public availability guarantees the reproducibility and accessibility of academic and applied research. Each dataset is further synthesized in Table 1, which includes its key features, advantages, as well as disadvantages.

Table 1. Key features, advantages, and disadvantages of selected dataset

Dataset	Key features	Advantage(s)	Disadvantage(s)
CEDAR (Srihari & Leedham, 2001)	55 signers, 24 genuine 24 forged samples each, total of 2640 grayscale images	Balanced and clean dataset, advisable for benchmarking	Limited signer diversity, English only, less suitable for cross-script research
BHSig260 (Kathuria, 2022)	260 signers (100 Bengali, 160 Hindi), 24 genuine 30 forged samples each	Large-scale, multilingual script support, high variation	Limited to Indic script, variability in forgery quality
ICDAR 2011 SigComp (Liwicki et al., 2011)	106 signers, 64 signers for training, 42 signers for testing,	Supports writer-independent evaluation	Dutch only, some outdated samples, limited scalability
Kaggle signature verification dataset (Reni, 2019)	2149 pre-processed Dutch signature images	Easy to use, curated format, ideal for fast prototyping	Medium sized, no strict train/test split, limited cultural diversity
Kaggle handwritten signatures of Divyansh Rai (Rai, 2020)	30 signers, 5 genuine 5 forged samples each, total of 300 offline handwriting samples	Lightweight, useful for small-scale or educational projects	Quite small for deep learning models, limited variability and generalization

3.2 Data augmentation techniques

Data augmentation is often crucial for improving the robustness and generalizability of deep learning models, particularly when working with limited datasets such as those common in offline signature forgery detection. Nevertheless, data augmentation approaches were predominantly absent from the examined studies. Traditional augmentation methods were reported in a restricted number of publications (Swamy et al., 2024), specifically applying geometric transformations such as image flipping, rotation, cropping, and contrast or brightness adjustments to artificially increase training diversity. Despite their simplicity, these methods may offer only marginal improvements in simulating skilled forgeries, which frequently encompass nuanced, high-fidelity handwriting characteristics that such techniques are incapable of replicating.

The lack of complete augmentation in most studies reveals a fundamental gap in the existing research landscape. This constraint has obvious implications for model overfitting and poor generalization, especially when evaluated against unknown or degraded samples. In response, this review proposes the integration of Generative Adversarial Networks (GANs) specifically Signature GANs, as a more advanced

data augmentation technique. GANs provide a unique benefit by generating realistic synthetic forgeries that closely resemble variances in human handwriting (Goodfellow et al., 2017; Wang & Jia, 2019). In contrast to traditional augmentations, GANs learn from authentic data distributions and generate intricate and realistic forgeries, which introduce controlled variability and improve the training set diversity. These properties allow GANs to surpass the constraints of traditional data augmentation techniques, which frequently fail to capture the intricate, skilled-level features of handwritten signatures. As a result, GAN-based augmentation is essential for mitigating dataset scarcity and improving model robustness in the offline signature forgery detection process. This augmentation strategy not only strengthens model learning but also enhances resilience against adversarial inputs and real-world noise, supporting the forensic readiness of future signature forgery detection systems.

3.3 Comparative analysis

This section provides a synthesis of the accuracy findings obtained from the total studies included in review (n=15) which are shown in Table 2. All studies employed the deep learning methods to the forgery detection of offline handwriting signatures. Accuracy is selected as the primary metric as it reflects as the most frequently reported, and interpretable metric. Accuracy measures the proportion of correctly classified signatures (genuine or forged) relative to the total number of samples.

Table 2. Accuracy comparative analysis

Study	Deep learning method	Dataset	Accuracy result (%)
Joe Harris & Anitha (2023)	Siamese network, Euclidean distance, contrastive loss function	CEDAR	100.0
Chokshi et al. (2023)	Siamese network, scattering wavelets learning approach, Euclidean distance	CEDAR	99.91
Balaji et al. (2024)	CNN with hyperparameter tuning	Kaggle signature verification dataset by RobinReni	99.8
Emberi et al. (2023)	Siamese network	Kaggle signature verification dataset by RobinReni	99.756
Majumder et al. (2023)	Siamese Transformer network, triplet loss function	CEDAR	99.17
Anitha et al. (2024)	Siamese network, Harris corner feature detection	BHSig260 (Hindi only)	98.9
Swamy et al. (2024)	Hybrid of autoencoders and CNN	CEDAR	98.4848
Shirisha et al. (2024)	VGG16 with transfer learning	Kaggle signature verification dataset by RobinReni	98.26
Chokshi et al. (2023)	Siamese network, scattering wavelets learning approach, Euclidean distance	ICDAR 2011 SigComp	97.76
Minh Tram & Chau (2024)	Pre-trained VGG16 with hyperparameter tuning	CEDAR	96.14
Reddy et al. (2024)	Siamese network, Euclidean distance, contrastive loss function	ICDAR 2011 SigComp	96.0
Krishna & Bhuvaneswari (2023)	20 hidden layers of multi-layer perceptron	Kaggle handwritten signatures by Divyansh Rai	92.4
Tehsin et al. (2024)	Siamese network, Euclidean distance, triplet loss function	BHSig260	91.5
Joe Harris & Anitha (2023)	Siamese network, Euclidean distance, contrastive loss function	BHSig260 (Hindi only)	87.13
Tehsin et al. (2024)	Triplet loss Siamese network	CEDAR	86.1
Tarek & Atia (2022)	Pre-trained ResNet50	Kaggle handwritten signatures by Divyansh Rai	82.0
Chaturvedi & Jain (2022)	Ensemble feature extractor and classifier, concatenation of geometrical features and features extracted from pre-trained MobileNet	BHSig260 (Hindi skilled forgery)	69.1
Jain et al. (2021)	Shallow Siamese network, Euclidean distance, contrastive loss function	Kaggle signature verification dataset by RobinReni	73.0 (Precision)

Table 2 reveals that offline signature forgery detection with deep learning model can be divided into three (3) categories: high, moderate, and low performances. Deep learning models that achieve more than 90% accuracy belongs under high performance category. Meanwhile, deep learning models with accuracy between 80% to 89% falls under the moderate performance category, and finally the models with accuracy below 80% is considered low performance deep learning models. The research conducted by Joe Harris & Anitha (2023) attained flawless accuracy of 100% utilizing a Siamese network with contrastive loss on the CEDAR dataset, ranking among the most effective deep learning model.

Siamese networks dominate the offline signature forgery detection field because of their robustness and adaptability across datasets. However, the choice of dataset has a significant impact on model performance, with CEDAR (Srihari & Leedham, 2001) and Kaggle signature verification dataset by RobinReni (Reni, 2019) provides higher accuracies, but datasets such as BHSig260 (Kathuria, 2022) particularly with skilled forgeries pose larger hurdles. This can be seen through the work by Chaturvedi & Jain (2022) where they reported only 69.1% accuracy despite employing a Siamese network, underscoring how dataset-specific challenges can affect results. The study by Jain et al. (2021) reported the precision results instead of accuracy. On the other hand, integrating advanced components such as scattering wavelets, Transformer blocks, or Autoencoders significantly boost forgery detection performance, suggesting captivating paths for future research.

4. DISCUSSIONS AND FUTURE WORK

This section discusses several deep learning models that achieve high performance (accuracy more than 90%) in forgery detection of offline handwriting signatures. Notably, Siamese networks, Siamese networks with scattering wavelets learning approach, Siamese Transformer network, VGG16 with transfer learning, and CNNs integrated with autoencoders have consistently achieved accuracy ranging from 96% to 100% across benchmark datasets. These findings highlight the potential of embedding-based similarity measures (Siamese networks), spatial encoding (Transformers), and convolutional feature extraction (VGG16 and CNNs) in learning offline signature features effectively.

The Siamese networks (Joe Harris & Anitha, 2023; Emberi et al., 2023; Reddy et al., 2024) dominate the landscape, often coupled with enhancements such as scattering wavelets learning approach (Chokshi et al., 2023) or Transformer network (Majumder et al., 2023). Siamese networks outperform the standard deep learning model i.e., standalone CNN, which often inadequately captures particular signature styles, resulting in suboptimal performance. Their success lies in their architecture's ability to learn relative similarity between signature pairs, allowing them to generalize well even with limited samples. Typically, a Siamese network consists of two identical sub-networks sharing weights and parameters, each receiving a different input image. The outputs are compared using a distance metric, usually Euclidean distance to determine its similarity. This architecture is particularly effective when trained with a contrastive loss function, which minimizes the distance between matching pairs while maximizing it for non-matching ones (Hadsell et al., 2006). Another alternative is the triplet loss function. This function compares an anchor image to both a positive and a negative sample and optimizes the network so that the anchor is closer to the positive than the negative by a specified margin (Schroff et al., 2015). While contrastive loss is computationally efficient with fewer samples, triplet loss frequently delivers greater discriminatory power in complex classification circumstances, especially when there are subtle differences in signatures. Table 3 depicts the differences between contrastive and triplet loss functions according to several criterion.

Table 3. Contrastive loss function vs triplet loss function

Criterion	Contrastive loss function	Triplet loss function	Recommended use case
Input structure	Pairwise input (genuine-genuine or genuine-forgery)	Triplet input (anchor, positive, negative)	Triplet loss is preferred when subtle intra-class differences are present
Loss objective	Minimize distance for similar pairs, maximize for dissimilar ones	Ensures anchor is closer to positive than to negative by a specified margin	Contrastive loss is effective for simpler verification tasks
Training efficiency	Faster training with fewer samples	Requires more structured triplet selection, leading to slower training	Use contrastive when data is limited or less diverse
Discriminative power	Moderate	Higher in complex decision boundaries	Triplet loss for fine-grained classification such as skilled forgery detection
Implementation complexity	Relatively straightforward	More complex due to triplet mining strategies	Contrastive for prototype verification; Triplet for high security systems

The comparisons in Table 3 suggest that in low-resource or low-variation settings, contrastive loss offers a simpler and effective approach. Meanwhile, for high-stakes applications involving skilled forgeries or cross-writer verification, triplet loss provides superior performance and granularity.

Furthermore, recent advancements in offline handwriting signatures forgery detection have investigated hybrid architectures that integrate the strengths of Siamese networks with Transformer-based encoders. This architecture employs a Siamese network, where each branch incorporates a Transformer encoder that analyzes signature embeddings via self-attention methods. In contrast to CNNs that concentrate mostly on local spatial attributes, Transformer-based modules allow the model to comprehend global structural linkages among strokes, curves, and spacing abnormalities, which are frequently essential for identifying competent forgeries. Table 4 shows the advantages and disadvantages of integrating the Transformer encoder into the Siamese network.

Table 4. Advantages and disadvantages of transformer encoder

Advantage(s)	Disadvantage(s)
Captures both local texture and global signature structure	Transformer layers add significant parameters thus, require careful model optimization
Reduces overfitting on writer specific traits	Self-attention mechanism typically performs better with larger training datasets
More adaptable to variable-length inputs	May suffer from convergence issues if not properly regularized or pre-trained

Meanwhile, the VGG16 deep learning model, a well-established architecture in image classification, also proves effective when adapted to offline handwriting signature forgery detection, particularly with transfer learning (Shirisha et al., 2024) and hyperparameter tuning (Minh Tram & Chau, 2024). VGG16's architecture is composed of 13 convolutional layers, followed by three fully connected layers, all utilizing small (3x3) convolutional filters with ReLU activations. This deep and uniform structure enables it to extract rich hierarchical features from input images, rendering it particularly well-suited for tasks such as signature verification, where subtle stroke differences are significant (Simonyan & Zisserman, 2014). Nevertheless, their efficacy is contingent upon the dataset and the fine-tuning strategy that is implemented. VGG16 is susceptible to overfitting and slow training when the dataset size is limited due to its extensive parameter set.

On the other hand, hybrid deep learning models that combine CNNs with autoencoders (Swamy et al., 2024) also demonstrate potential by compressing signature representations and reconstructing them to emphasize distinguishing features. Autoencoders are unsupervised learning models that acquire the ability to encode input data into a lower-dimensional latent space and subsequently decode it back to its original form (Khan et al., 2020). Within the context of CNNs with autoencoders, the CNN acts as a powerful feature extractor, while the autoencoder refines these features by reconstructing the signature image. This reconstruction process forces the hybrid model to preserve the most salient and distinguishing features of a signature, effectively reducing noise and irrelevant variations. As a result, this model becomes more robust in distinguishing between genuine and forged signatures, particularly in datasets with limited labeled samples.

However, despite the encouraging performance in benchmark scenarios, three major challenges persist that impact real-world deployment:

- a) **Dataset challenges.** Most studies relied on limited or specific datasets, which reduces the diversity of signature variations in terms of culture, language, and writing instruments. Subsequently, models trained on these datasets tend to show high performance in controlled environments but struggle to generalize. The implication is a significant gap in model robustness when deployed in practical and diverse contexts.
- b) **Discriminative model challenges.** While discriminative models like Siamese and VGG-based architectures achieve high accuracy, they often operate as black boxes. Their dependency on dataset-specific features makes them prone to overfitting, and the lack of transparency reduces trust in high-stakes environments like banking or legal authentication. Moreover, the absence of interpretability tools limits understanding of model decisions.
- c) **Deployment challenges.** Most of the developed models lack resilience to noise and real-world variability. The accuracy scores reported in the literature assume clean and well-aligned signature samples, while actual verification systems must contend with occlusions, blurred scans, and variable lighting. As a result, model robustness under real-world deployment remains an open challenge.

4.1 Towards end-to-end generative-aware conceptual framework

An end-to-end generative-aware conceptual framework is proposed. It unifies several essential innovations: data augmentation, hybrid architecture, noise resilience, and explainable verification. This proposed framework addresses current limitations in dataset diversity, model generalization, real-world robustness, and interpretability. It is further organized into four interconnected stages as shown on Fig. 2.

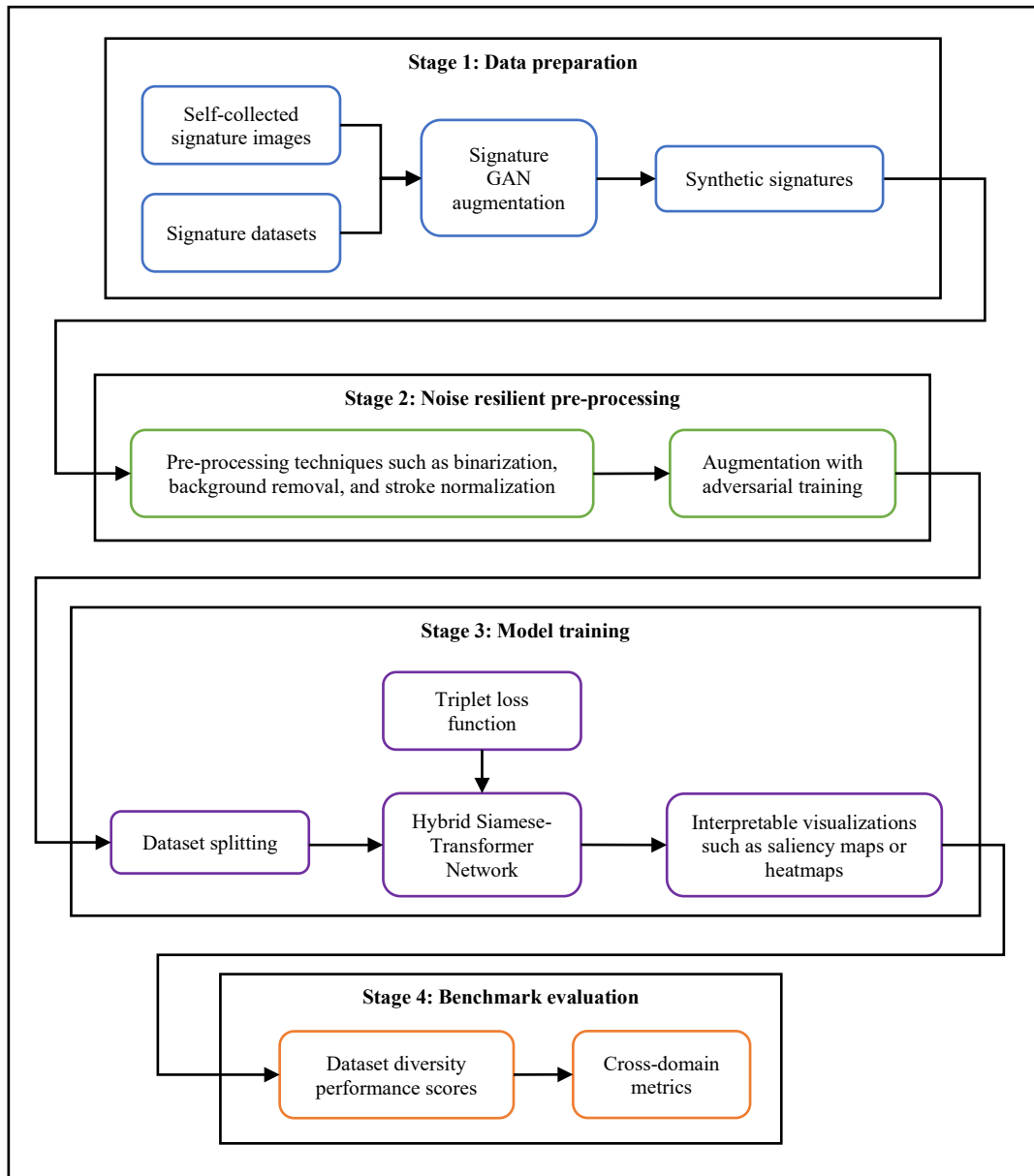


Fig. 2. Proposed end-to-end generative-aware conceptual framework

Stage 1: Data Preparation. Self-collected signature images as well as signature datasets such as CEDAR (Srihari & Leedham, 2001) and BHSig260 (Kathuria, 2022) are composed as foundational data. A specialized Signature Generative Adversarial Network (GAN) is trained to generate synthetic forgeries that simulate variations in signatures' style and structure. It utilizes Generative Adversarial Networks (GANs) to synthetically produce realistic signature samples by learning the intricate variations found in handwriting, such as pressure, stroke curvature, and pen lifts. Introduced by Goodfellow et al. (2017), GANs consist of two core components: a generator that aims to produce synthetic data samples resembling real data, and a discriminator that evaluates whether a given sample is real or generated. These two networks are trained

simultaneously in a minimax game, where the generator strives to deceive the discriminator, and the discriminator attempts to accurately identify genuine versus forged samples. This adversarial training process allows the Signature GAN to model complex data distributions more effectively than conventional augmentation methods, such as rotation, flipping, and scaling. The visual process of GANs is shown in Fig. 3.

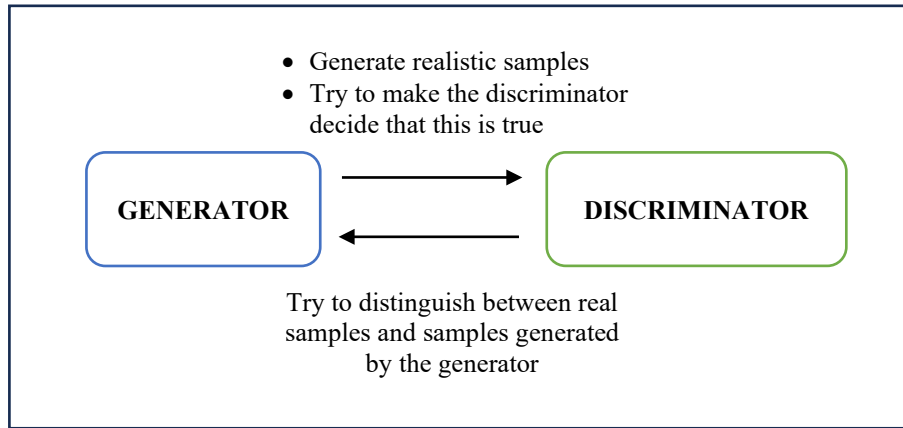


Fig. 3. GANs process between a generator and a discriminator

From a computational standpoint, while GANs require extended training durations and meticulous tuning to ensure stability, they significantly reduce the necessity for repetitive, hand-crafted augmentation techniques. This optimizes the data enrichment procedure and reduces design intricacy. Researchers like Dash et al. (2023) have emphasized the significance of employing GANs in tasks where data authenticity and variety are critical, such as signature forgery detection. The use of Signature GANs, in particular enable the generation of scalable and balanced datasets that encompass diverse forgery scenarios, eliminating the necessity for manual intervention or additional data collection. This technique efficiently mitigates data scarcity while incorporating controlled variability during training, essential for creating robust and generalizable forgery detection model.

Stage 2: Noise resilient pre-processing. As the real-world variability involve noises such as occlusions, blurred scans, and variable lighting, this stage ensures the dataset quality through preprocessing steps such as binarization (converting grayscale or color signature images into binary black-and-white format to enhance contrast), background removal (eliminating unnecessary background textures or lines that may interfere with signature contours), and stroke normalization (adjusting stroke width, continuity, and alignment to standardize writing patterns across samples). It is further strengthened with adversarial training to increase robustness and resilience to unexpected input degradation.

Stage 3: Model training. The core architecture is a Hybrid Siamese-Transformer Network, which merges the strengths of Siamese structures and Transformer-based encoders. The Siamese network allows for effective use of triplet loss to measure similarity between anchor-positive-negative pairs, while the Transformer encoder captures long-range dependencies in the stroke sequence. In addition, an auxiliary binary classification submodule is introduced at the end of this hybrid architecture, so that it can strongly improve the detection between genuine and forged signatures, particularly in skilled forgery detection. Interpretable visual outputs such as saliency maps and heatmaps are added to contribute to the system's transparency and trustworthiness.

Stage 4: Benchmark evaluation. This final stage aims to assess the model's generalizability and robustness across diverse real-world conditions. This stage consists of two key components: dataset diversity performance scores and cross-domain metrics. The first component evaluates how well the model

performs across multiple signature datasets with varying languages, writing styles, cultural characteristics, and acquisition conditions. Metrics such as accuracy, precision, and recall are used to determine the model's adaptability and resistance to overfitting. The second component, cross-domain metrics, focuses on how well the model transfers to unseen domains such as new user populations or degraded documents. Together, these metrics ensure the model is not only benchmark-strong but also practically deployable, robust, and fair in real-world applications.

5. CONCLUSIONS

This review paper provided an extensive examination of AI-driven offline handwriting signature forgery detection, focusing specifically on deep learning and novel generative adversarial networks methodologies. A rigorous analysis of existing models, datasets, and performance indicators revealed that Siamese-based architecture prevails in the field due to their capacity to learn relative similarity between pairs of signatures. Although great accuracy has been attained on controlled datasets like CEDAR (Srihari & Leedham, 2001), deployment issues remain in real-world settings because of differences in dataset complexity, noise, and extraneous fluctuations. Meanwhile, hybrid Siamese-Transformer models and Generative-Aware frameworks exhibit considerable potential for addressing existing constraints, particularly by improving contextual learning and resilience while mitigating overfitting concerns.

Furthermore, the role of data augmentation is equally significant, since the integration of Signature GANs produces realistic, high-fidelity synthetic samples that mitigate the ongoing issue of data scarcity. In contrast to conventional augmentation techniques that utilize fixed, manually designed transformations, Signature GANs acquire intricate handwriting distributions directly from authentic examples. This facilitates the creation of sophisticated forgeries that more accurately represent real-world variances. Thus, this augmentation technique not only improves training diversity but also strengthens model robustness under degraded or adversarial conditions.

Future research should prioritize the integration of generative adversarial networks for both data augmentation and forgery detection, development of noise-tolerant models, and embedding explainable AI mechanisms to ensure transparency and trustworthiness in forensic applications. By solving these crucial shortcomings, AI-driven offline handwriting signature forgery detection can move closer to reaching forensic-grade reliability suited for implementation in high-risk legal, financial, and security situations.

6. ACKNOWLEDGEMENTS/FUNDING

This research was not funded by any grant.

7. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

8. AUTHORS' CONTRIBUTIONS

Safura Adeela Sukiman: Conceptualisation, methodology, formal analysis, investigation, data curation, writing-original draft, and visualization; **Nor Azura Husin:** Conceptualisation, methodology, validation, resources, writing-review and editing, and supervision; **Hazlina Hamdan:** Conceptualisation, validation, resources, and supervision; **Masrah Azrifah Azmi Murad:** Conceptualisation, validation, resources, and supervision.

9. REFERENCES

- Anitha, A., Priya, M., Kamaraj, B., Yadav, N. K., & Srivastav, S. (2024). Handwritten signature forgery identification and prediction using Harris corner detection and Siamese network. In *2024 Second International Conference on Emerging Trends in Information Technology and Engineering (ICETITE)* (pp. 1-5). IEEE. <https://doi.org/10.1109/ic-etite58242.2024.10493595>
- Balaji, Y., Vikas Kedar, R., Virupakshi, H., & Anandan, P. (2024). Crafting trust with convolutional neural networks and hyperparameter tuning for precision signature verification in insurance claim. In *2024 International Conference on Advances in Data Engineering and Intelligent Computing Systems (ADICS)* (pp. 1-6). IEEE. <https://doi.org/10.1109/adics58448.2024.10533503>
- Chaturvedi, P., & Jain, A. (2022). Feature ensemble based method for verification of offline signature images. In *2022 International Conference on Machine Learning, Big Data, Cloud and Parallel Computing (COM-IT-CON)* (pp. 710-714). IEEE. <https://doi.org/10.1109/com-it-con54601.2022.9850628>
- Chokshi, A., Jain, V., Bhope, R., & Dhage, S. (2023). SigScatNet: A siamese + scattering based deep learning approach for signature forgery detection and similarity assessment. In *2023 International Conference on Modeling, Simulation & Intelligent Computing (MoSICom)* (pp. 480-485). IEEE. <https://doi.org/10.1109/mosicom59118.2023.10458765>
- Dash, R., Bag, M., Pattanayak, D., Mohanty, A., & Dash, I. (2023). Automated signature inspection and forgery detection utilizing VGG-16: A deep convolutional neural network. In *2023 2nd International Conference on Ambient Intelligence in Health Care (ICAHC)* (pp. 01-06). IEEE. <https://doi.org/10.1109/icaihc59020.2023.10431430>
- Emberi, N. B., Mohan, A., Naphade, C. A., & Ransing, R. (2023). Harnessing deep neural networks for accurate offline signature forgery detection. In *2023 7th International Conference on Intelligent Computing and Control Systems (ICICCS)* (pp. 619-626). <https://doi.org/10.1109/iciccs56967.2023.10142593>
- Engin, D., Kantarci, A., Arslan, S., & Ekenel, H. K. (2020). Offline signature verification on real-world documents. In *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)* (pp 806-808). <https://doi.org/10.1109/cvprw50498.2020.00412>
- Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., & Bengio, Y. (2017). GAN (Generative Adversarial Nets). *Journal of Japan Society for Fuzzy Theory and Intelligent Informatics*, 29(5), 177. https://doi.org/10.3156/jsoft.29.5_177_2
- Hadsell, R., Chopra, S., & LeCun, Y. (2006). Dimensionality reduction by learning an invariant mapping. *2006 IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, 2, 1735–1742. <https://doi.org/10.1109/cvpr.2006.100>
- Hafemann, L. G., Sabourin, R., & Oliveira, L. S. (2017). Offline handwritten signature verification - Literature review. In *2017 Seventh International Conference on Image Processing Theory, Tools and Applications (IPTA)* (pp. 1-8). <https://doi.org/10.1109/ipta.2017.8310112>
- Jain, S., Khanna, M., & Singh, A. (2021). Comparison among different CNN architectures for signature forgery detection using Siamese neural network. In *2021 International Conference on Computing, Communication, and Intelligent Systems (ICCCIS)* (pp. 481–486). <https://doi.org/10.1109/icccis51004.2021.9397114>
- Joe Harris, S. K., & Anitha, J. (2023). An improved signature forgery detection using modified CNN in siamese network. In *2023 Second International Conference on Electrical, Electronics, Information and Communication Technologies (ICEEICT)* (pp. 1–5).

<https://doi.org/10.1109/iceeict56924.2023.10156951>

- Kathuria, I. (2022). *Handwritten signature datasets* [Data set]. Kaggle. <https://www.kaggle.com/datasets/ishanikathuria/handwritten-signature-datasets>.
- Khan, A., Sohail, A., Zahoora, U., & Qureshi, A. S. (2020). A survey of the recent architectures of deep convolutional Neural Networks. *Artificial Intelligence Review*, 53(8), 5455–5516. <https://doi.org/10.1007/s10462-020-09825-6>
- Krishna, C. A., & Bhuvaneswari, R. (2023). Offline signature forgery detection using Multi-Layer Perceptron. In *2023 3rd Asian Conference on Innovation in Technology (ASIANCON)* (pp. 1-4). <https://doi.org/10.1109/asiancon58793.2023.10269874>
- Liwicki, M., Blumenstein, M., van den Heuvel, E., Berger, C. E. H., Stoel, R. D., Found, B., Chen, X., & Malik, M. I. (2011). *Hugging Face* [Data set]. ICDAR 2011 Signature Verification Competition (SigComp2011). <https://huggingface.co/datasets/laurent/ICDAR-2011>
- Majumder, P., Joaa, A. M., Rahman Rhythm, E., Kabir Mehedi, M. H., & Alim Rasel, A. (2023). Siamese-transformer network for offline handwritten signature verification using few-shot. In *2023 26th International Conference on Computer and Information Technology (ICCIT)* (pp. 1–6). IEEE. <https://doi.org/10.1109/iccit60459.2023.10441035>
- Minh Tram, P. H., & Chau, D. N. (2024). Offline signature forgery detection using image processing. In *2024 13th International Conference on Control, Automation and Information Sciences (ICCAIS)* (pp. 1-6). IEEE. <https://doi.org/10.1109/iccais63750.2024.10814324>
- Page, M. J., McKenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann, T. C., Mulrow, C. D., ... & Moher, D. (2021). The PRISMA 2020 statement: an updated guideline for reporting systematic reviews. *BMJ*, 372.
- Rai, D. (2020). *Handwritten signatures dataset* [Data set]. Kaggle. <https://www.kaggle.com/datasets/divyanshrai/handwritten-signatures>
- Reddy, M. S., Lakshmi, A. A., Reddy, G. S., Madhavi, B. K., Panigrahi, B. S., & Mohan, V. (2024). Signature forgery detection using siamese-convolutional neural network. In *2024 1st International Conference on Cognitive, Green and Ubiquitous Computing (IC-CGU)* (pp. 1–5). IEEE. <https://doi.org/10.1109/ic-cgu58078.2024.10530761>
- Reni, R. (2019). *Signature verification dataset* [Data set]. Kaggle. <https://www.kaggle.com/datasets/robinreni/signature-verification-dataset>.
- Schroff, F., Kalenichenko, D., & Philbin, J. (2015). FaceNet: A unified embedding for face recognition and clustering. In *2015 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)* (pp. 815–823). IEEE. <https://doi.org/10.1109/cvpr.2015.7298682>
- Shirisha, N., Pranitha, V. P., Balam, A., Kalpana Chowdary, M., & Shekar, K. (2024). A learning-based intelligent method for signature forgery detection using pre-trained deep models. In *2024 Second International Conference on Intelligent Cyber Physical Systems and Internet of Things (ICoICI)* (pp. 1337-1345). IEEE. <https://doi.org/10.1109/icoici62503.2024.10696616>
- Simonyan, K., & Zisserman, A. (2014). Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556*.
- Soelistio, E. A., Hananto Kusumo, R. E., Martan, Z. V., & Irwansyah, E. (2021). A review of signature recognition using machine learning. In *2021 1st International Conference on Computer Science and Artificial Intelligence (ICCSAI)* (pp. 219-223). <https://doi.org/10.1109/iccsai53272.2021.9609732>

- Srihari, S. N., & Leedham, G. (2001). *CEDAR signature verification dataset* [Data set]. Center of Excellence for Document Analysis and Recognition (CEDAR), University at Buffalo. <http://www.cedar.buffalo.edu/NIJ/>
- Swamy, B. N., Kumar, D. L., Jyothi, K. L., Harika, M. V., Kancharla, G., & Karthik, B. R. (2024). Offline signature verification with AUTOENCODERCNN hybrid feature extraction for improved fake signature detection. In *2024 International Conference on Social and Sustainable Innovations in Technology and Engineering (SASI-ITE)* (pp. 418–423). IEEE. <https://doi.org/10.1109/sasi-ite58663.2024.00085>
- Tarek, O., & Atia, A. (2022). Forensic handwritten signature identification using Deep Learning. In *2022 IEEE 9th International Conference on Sciences of Electronics, Technologies of Information and Telecommunications (SETIT)* (pp. 185–190). IEEE. <https://doi.org/10.1109/setit54465.2022.9875697>
- Tehsin, S., Hassan, A., Riaz, F., Nasir, I. M., Fitriyani, N. L., & Syafrudin, M. (2024). Enhancing signature verification using triplet Siamese similarity networks in digital documents. *Mathematics*, 12(17), 2757. <https://doi.org/10.3390/math12172757>
- Wang, S., & Jia, S. (2019). Signature handwriting identification based on generative adversarial networks. *Journal of Physics: Conference Series*, 1187(4), 042047. <https://doi.org/10.1088/1742-6596/1187/4/042047>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

CNN Integrated Mobile Application: Food Image Recognition for Recipe Generation

Muhammad Imran Nor Azlan Shah¹, Norlina Mod Sabri^{2*}, Gloria Jennis Tan³,
Zhiping Zhang⁴

¹Zen Computer System, Persiaran Flora, Cyber 12, 63000 Cyberjaya, Selangor, Malaysia.

^{2,3}Faculty of Computer and Mathematical Sciences, UiTM Terengganu Branch, Kuala Terengganu Campus, 21080 Kuala Terengganu, Terengganu, Malaysia.

⁴College of Computer and Mathematics, Xinyu University, Jiangxi, P. R. China.

ARTICLE INFO

Article history:

Received 27 June 2025

Revised 31 July 2025

Accepted 6 August 2025

Published 1 September 2025

Keywords:

CNN

Mobile Application

Food Image

Recognition

DOI:

10.24191/jcrinn.v10i2.554

ABSTRACT

The rapidly advancing of the digital world has encouraged the use of mobile applications in almost every aspects of our everyday life. This includes transforming the way we obtain our meal, whether to order from food providers or simply cook for ourselves. The CNN Integrated Mobile Application for Food Recipe Generation is intended to improve users' culinary experiences by offering suggestions for recipes and intelligent ingredient recognition. This research explores the Convolutional Neural Networks (CNN) algorithm to tackle the problems of effective ingredient management and minimize food waste through smartphone application. Through the mobile application, the ingredient photos can be scanned by users or uploaded, after which the CNN model processes the images to precisely identify the ingredients. The application provides users with a wide range of meal alternatives that are customized to their available ingredients by retrieving relevant recipes from external databases such as the Spoonacular API based on the recognized ingredients. The research methodology consists of 3 main phases, which are Data Preprocessing, CNN Implementation and Performance Evaluation. In this research, the CNN algorithm has generated a good and acceptable performance with more than 96% accuracy. This research has shown how machine learning, mobile development, and user-centric design can be successfully combined to create a useful tool for contemporary culinary demands. The app encourages a move towards more sustainable and thoughtful eating habits by acting as an incentive for change at the community level. When communities embrace these ideas, the app plays a key role in tackling more significant social issues associated with food waste, supporting international initiatives that are detailed in the UN Sustainable Development Goals (SDGs) for a more sustainable and responsible society.

^{2*} Corresponding author. E-mail address: norli097@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.554>

1. INTRODUCTION

In today's fast-paced society, finding a balance between quick meal preparation and making healthy food choices can sometimes be challenging. The evolving urban lifestyle has increased the demand for solutions that streamline our daily food decisions, while ensuring a balance between efficiency and mindful health practices. Traditional methods, including manual seeking and searching is not sufficient anymore, thus requires a new revolutionary method in providing ease while actively seeking for good food. Furthermore, people nowadays frequently consume food prepared by third-party entities such as restaurants, catering services, and takeaways, often through online delivery apps for added convenience (Chopra et al., 2023). As a result, looking for cooking recipes and ingredient information seems to be a tedious task. It might be laborious and time-consuming for people to manually check the ingredients that are available, choose what to prepare and look up appropriate recipes especially for those with hectic schedules. Due to the lack of proper planning, consumers may find it difficult to use products before their expiration dates, resulting in a greater chance that they would be wasted. This scenario results in financial losses and food waste (Rodrigues et al., 2023). A more effective and technologically advanced solution that modernizes the process of ingredient identification and recipe recommendation is becoming more and more necessary considering these disadvantages (Dsouza et al., 2024).

Currently, recipe recommendation systems have become widespread in order to provide clients with customized recipe recommendations based on their dietary preferences, level of cooking proficiency, and availability of ingredients. However, the majority of recipe recommendation systems rely on textual inputs or predefined categories of ingredients, which might not accurately reflect what the user has in their pantry (Anitha et al., 2023). Based on these motivations, a mobile solution is proposed with the ultimate aim of enhancing the users' quality of life by providing a seamless and intelligent food recipe assistance. Recent studies suggest that interactive and personalized tools, such as mobile applications, can be more effective in supporting real-time decision-making (Kansaksiri et al., 2023). The proposed solution is to overcome the issues of global food loss, hectic schedules, and inadequate recipe recommendations by utilizing CNN algorithm. The objective of the research is to develop and explore the Convolutional Neural Network (CNN) algorithm for the mobile application in the food image recognition. This study proposes the implementation of CNN to revolutionize the identification of food ingredients, initially focusing on 5 classes of ingredients including egg, lettuce, tomato, carrot and onion. In the context of the food scanning app, image recognition technology serves as the backbone for accurate and efficient identification of ingredients from user-provided images. The selection of CNN algorithms is a strategic one, driven by their robustness in handling image data and their capacity for scalability. CNN is a cutting-edge technique that is frequently used to train computers to recognize and classify images to obtain a trained dataset (Anitha et al., 2023).

In this study, the mobile application will recommend recipes based on the image detected. This will make finding recipes quicker and simpler than the manual searching. The development of the application supports the proposed issues by using image recognition to swiftly identify ingredients, addressing the consequences of poor planning and over-preparation. It encourages home cooking by recommending recipes based on available ingredients, aligning with modern lifestyles. Utilizing CNNs capability, the mobile application will revolutionize how individuals manage food, offering a sustainable, time-efficient, and personalized solution to daily culinary challenges. This paper is arranged into 5 main sections, beginning with the Introduction section. The following sections cover the Literature Review, Methodology, Results and Discussion and finally the Conclusion and Recommendation.

2. LITERATURE REVIEW

This section presents brief reviews about similar works of the research and CNN algorithm. In similar works, several latest deep learning related research projects are compiled and summarized. The review on CNN discussing its advantages in food image recognitions.

2.1 Similar works

This section presents the deep learning research projects which have involved food identification and classification. A research was initiated in China to curb obesity, by using Yolo to recognize fast food and display its information (Chen & Chiang, 2025). Another research has identified real time food items using YOLOv5 and also display its information (Anand et al., 2024). Both research have showed improved performance of yolo in their projects. In UK, in order to reduce wastage, a food image recognition research has been conducted using CNN and the accuracy has been improved (Boyd et al., 2024).

A recipe recommendation system using YOLOv5 for ingredient detection and calorie calculation is proposed by Swain et al. (2023). Meanwhile, a similar work introduces a system that employs Convolutional Neural Networks (CNN) for personalized recipe recommendations based on user-provided ingredient images (Anitha et al., 2023). Both approaches offer their own strengths as the system using YOLOv5 includes consideration of nutritional value while CNN focuses on user experience and tailored recipe recommendations. CNN also addresses ingredient scarcity, simplifying the recipe suggestion process, and enhancing user experience. However, both methods also come with similar weaknesses in which accuracy of ingredient detection is highly dependent on image quality and data training.

In the research by Li (2023), the article proposes a system using CNN and Transformer for predicting food recipes from images. Strengths lie in CNN's feature extraction, but weaknesses include difficulty capturing fine details. Combining CNN and Transformer models can overcome these weaknesses, improving recipe prediction accuracy, and suggesting that knowledge distillation enhances model portability. Chopra et al. (2023) presents a solution for image-recipe association in Indian cuisine using machine learning. Strengths include accurate predictions for a specific cuisine. A potential weakness is the limitation to Indian cuisine, potentially hindering the retrieval of recipes from other cuisines.

Based on literatures, this research has been motivated by the performance of CNN in the previous food image recognition research projects. CNN has proven to be able to produce good performance in the food detection and classification problems. By linking input data to spatial features, the CNN algorithm can detect and extract key information from images, making it highly useful across various applications (Sri & Balakrishnan, 2024). Table 1 shows the similar works using deep learning algorithms and their results in the food image recognitions.

Table 1. Similar works using deep learning in food image recognitions

No	Title	Problem	Result	Reference
1	Deep learning-based automatic food identification with numeric label	Food identification and calorie management system is proposed due to obesity issue	Yolo classify and count objects within improved time complexity	Chen and Chiang (2025)
2	Food Recognition Using Deep Learning for Recipe and Restaurant Recommendation	Identify various food items in real time from an image or video of multi-object meals	YOLOV5 obtained average accuracy over 80%	Anand et al. (2024)

3	Fine-Grained Food Image Recognition: A Study on Optimising Convolutional Neural Networks for Improved Performance	Increasing issue of food waste	Improved accuracy for DenseNet	Boyd et al. (2024)
4	Recipe Recommendation Using Image Classification	Recipe recommendation systems use textual inputs or pre-defined categories of ingredients, not accurately reflect what the user has in their pantry	The use of CNNs provides high accuracy in identifying ingredients	Anitha et al. (2023)
5	Ingredients to Recipe: A YOLO-based Object Detector and Recommendation System via Clustering Approach	Website lack real time analysis of ingredients used for cooking	YOLO is a lot faster algorithm than its competitors, with a maximum frame rate of 45 FPS.	Swain et al. (2023)
6	Ingredient Detection, Title, and Recipe Retrieval from Food Images	Simply looking at a picture of food on social media does not provide access to the cooking process.	Recipe generation model presents notable advancements over prior models in the field	Chopra et al. (2023)
7	Predict Food Recipe from Image Using CNN/Transformer & Knowledge Distillation for Portability	Chef cannot remember the detailed steps to all the dishes, and ordinary citizens do not have the time or energy to create their recipes.	The resulting model's complexity was minimized	Li (2023)

2.2 Convolutional neural network (CNN)

CNNs are an effective method for image classification in computer vision (Anitha et al., 2023). Unlike traditional image recognition algorithms, CNN excels in capturing intricate patterns and features within images, making them highly effective for identifying nuanced details in food ingredients. The hierarchical structure of CNN enables them to understand spatial relationships and contextual information, enhancing the precision of ingredient recognition. While CNN may be slower due to their deep architecture, the trade-off is increased accuracy, a crucial factor in the context of the app where precision is paramount. In contrast, certain image recognition algorithms might not have the intricacy and depth required to precisely identify minute details in food photos. Certain algorithms might process information more quickly, but at the expense of accuracy, which could result in incorrect recommendations and misclassifications. This comparison highlights CNN's applicability to the complex issue of culinary ingredient recognition. Classification of images using CNN is an optimal solution as its built-in convolutional layers reduce the high dimensionality of the image without losing its information (Navaprakash et al., 2025; Singh & Susan, 2023).

However, it's important to recognise CNN's shortcomings, particularly its high computational overhead, which causes processing times to increase. In spite of this limitation, the precision obtained by CNN is in perfect harmony with the app's goals of offering accurate ingredient identification and customised recipe suggestions. CNN's integration guarantees that accuracy comes before speed, a calculated move that improves both the app's general use and its ability to reduce food waste. Convolutional Neural Networks are very accurate networks for image classification and recognition in a short time (Kurian & Jacob, 2023). In conclusion, the food scanning app made a calculated decision to include the CNN algorithm because of its capacity to improve ingredient recognition and suggestion accuracy. The contrast with other algorithms highlights CNN's distinct benefits and establishes them as the best option for the app's complex image recognition tasks. Deep learning, especially convolutional neural networks (CNNs), has become the cornerstone of modern food image recognition systems. These models excel at automatically extracting features from images, enabling accurate classification and segmentation of food items (Gautam & Khanna, 2024).

3. METHODOLOGY

The research methodology consists of 3 main phases, which are Data Preprocessing, CNN Implementation and Performance Evaluation phase.

3.1 Data preprocessing

The dataset for this research has been obtained from Roboflow.com, a well-known online website used for storing and sharing datasets. The dataset used in this study was specifically selected to align with the research focus on ingredient identification and was sourced from Roboflow's extensive library. Roboflow contains the required datasets that fill in the criteria and requirements needed to develop the app, mainly the availability of datasets of the 5 chosen classes; which are egg, tomato, lettuce, onion and carrot.

Resizing the images is the first step in the preprocessing stage of creating the CNN model. By giving the CNN a constant input size, uniform sizing of all the images in the dataset improves its performance. It is essential to resize every image to a predetermined size before training in order to expedite this procedure and prevent any problems. Since they lighten the computational load and speed up processing, smaller image sizes are especially advantageous for mobile applications with low processing power which improves user experience overall. Image resizing is integrated into the code of CNN model development through the utilization of the `tf.keras.preprocessing.image_dataset_from_directory` function. The image size is set to 150x150 pixels, which is common for CNN.

3.2 CNN implementation

In this research project, the mobile application employs the CNN algorithm to precisely detect and categorize various food products, specifically eggs, tomatoes, lettuce, carrots, and onions. The CNN model was put into practice using the TensorFlow and Keras libraries. The architecture makes use of many convolutional layers, max-pooling layers, and fully linked layers to extract and learn features from images. Convolutional layers extract features, max-pooling layers minimize spatial dimensions, and fully connected layers perform the final classification. Then, the model is compiled using the Adam optimizer, employing categorical cross-entropy as the loss function, and accuracy is chosen as the evaluation metric. Fig. 1 shows the code snippet for building the CNN model.

```
# Define the model
model = Sequential([
    Conv2D(32, (3, 3), activation='relu', input_shape=(IMG_HEIGHT, IMG_WIDTH, 3)),
    MaxPooling2D(2, 2),
    Conv2D(64, (3, 3), activation='relu'),
    MaxPooling2D(2, 2),
    Conv2D(128, (3, 3), activation='relu'),
    MaxPooling2D(2, 2),
    Flatten(),
    Dense(512, activation='relu'),
    Dropout(0.5),
    Dense(5, activation='softmax') # Assuming 5 classes: egg, tomato, lettuce, carrot, onion
])

model.compile(optimizer='adam', loss='categorical_crossentropy', metrics=['accuracy'])
```

Fig.1. Code snippet of building the CNN model

Fig. 1 shows the CNN model is built using the Sequential class, which creates a linear stack for layer-by-layer definition. The CNN uses multiple layers to learn and extract information from input images. Convolutional layers extract features, max-pooling layers minimize spatial dimensions, and fully connected

layers perform the final classification. The first Conv2D layer uses 32 filters, each with a 3x3 kernel size, and the second layer uses 64 filters, each with a 3x3 kernel size. To incorporate nonlinearity into the network and allow the model to learn complicated patterns, the ReLU (Rectified Linear Unit) activation function is used. The model's input shape is defined as (150, 150, 3), indicating that the input images have a 150x150 pixel resolution and three color channels (RGB).

The MaxPooling2D layers are used to downsample the feature maps by selecting the maximum value inside each 2x2 zone, reducing spatial dimensions while retaining the most important information. The Flatten layer, which comes after the pooling layers, transforms the 2D feature maps into 1D vectors that may be used as input for the fully connected layers. The ultimate classification procedure is carried out by dense layers, sometimes referred to as fully connected layers. With 512 neurons, the initial Dense layer is a strong layer for feature learning. The output Dense layer uses the softmax activation function to create a probability distribution over the five food classes, which are represented by the five neurons in the layer. Finally, the model is compiled using the Adam optimizer, employing categorical cross-entropy as the loss function, and accuracy is chosen as the evaluation metric.

Processed images and labels are used to train the model. During the training process, the model is fed training data, weights are adjusted by backpropagation, and model parameters are improved. Using the dataset supplied by the train_generator, the model is trained for a predefined number of epochs via the model.fit() method. The model receives batches of training data from the train_generator, a data generator. In order to guarantee that the model is trained on different portions of the dataset in each epoch, it creates batches of images and labels. The model's performance is evaluated using validation data at the end of each training period, which enables tracking of the model's evolution and overfitting detection. Finally the model is converted to TensorFlow Lite in order to be used on mobile devices. A suite of tools called TensorFlow Lite makes it easier to run TensorFlow models on embedded, mobile, and Internet of things devices.

3.3 Performance evaluation

Evaluating the performance of the classification algorithm involves a few different metrics to measure the algorithm's effectiveness and accuracy. Accuracy, precision, recall and F1 are some of the metrics in performance measurements. The classified data obtained via the proposed classification algorithm would be measured in a confusion matrix. The confusion matrix contains the classifier's measurement parameters with four essential properties. Fig. 2 shows the confusion matrix.

		Actual Values	
		Positive (1)	Negative (0)
Predicted Values	Positive (1)	TP	FP
	Negative (0)	FN	TN

Fig. 2. Confusion matrix

Based on Fig. 2, TP (True positive) represents the total number of cases that are predicted true and correctly predicted, TN (True negative) represents the total number of cases that are predicted false and correctly predicted, FP (False positive) is the total number of cases that are predicted true but a wrong prediction, while FN (False negative) represents the total number of cases that are predicted false and is a wrong prediction. Using the values obtained via the confusion matrix, the performance evaluation metrics could easily be calculated using the formula in Eq. (1) – (4). Eq. (1) represents the measurement of accuracy, while Eq. (2) represents the measurement of precision. Eq. (3) represent the recall evaluation and Eq. (4) represents the F1 score evaluation.

$$\frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

$$\frac{TP}{TP + FP} \quad (2)$$

$$\frac{TP}{TP + FN} \quad (3)$$

$$2 * \frac{accuracy * recall}{accuracy + recall} \quad (4)$$

3.4 System architecture

System architecture shows the multiple layers between users, devices, networks, and data. Fig. 3 shows the system architecture of the project. The system architecture of the CNN Integrated Mobile Application for Food Recipe Assistance is meticulously planned to provide accurate and efficient ingredient recognition and recipe recommendations. The primary application of the system for processing photos and recognizing ingredients is the Convolutional Neural Network (CNN) algorithm. To begin the process, a sizable dataset consisting of images of five distinct ingredients was obtained from Roboflow. This dataset goes through a data preprocessing step where photos are shrunk and standardized in order to preserve consistency and enhance the model's performance. The CNN algorithm is then fed the processed data. The model is retained for later usage after being assembled with the proper configurations during the training process. Rigorous testing is conducted to assess the model's accuracy and ensure it meets the required performance standards. This design makes extensive use of the user interface (UI), which provides a clear-cut and easy-to-use experience. The user interface (UI) enables users to submit photos from their gallery and uses the device's camera to take shots of ingredients. Moreover, users can manually look up recipes by utilizing the UI's input interface. The CNN model examines the image once an ingredient has been scanned or uploaded in order to identify the item. Utilizing this recognition, the system retrieves relevant recipe recommendations from an external API, the Spoonacular API, which offers a large recipe database using this recognition. The user interface then displays the recommended recipes to the user, guaranteeing a smooth and interesting experience. The user is then presented with the suggested recipes via the UI, ensuring a seamless and engaging experience. Additionally, the integration of TensorFlow Lite optimizes the CNN model for mobile devices, providing quick and responsive ingredient recognition. The project's comprehensive approach is demonstrated by the system architecture, which ensures both technical efficiency and user satisfaction through the clear delineation of components and their interconnections. This architecture's application

demonstrates how cutting-edge technologies can be used to improve routine chores, making it a useful tool for both home cooks and food connoisseurs.

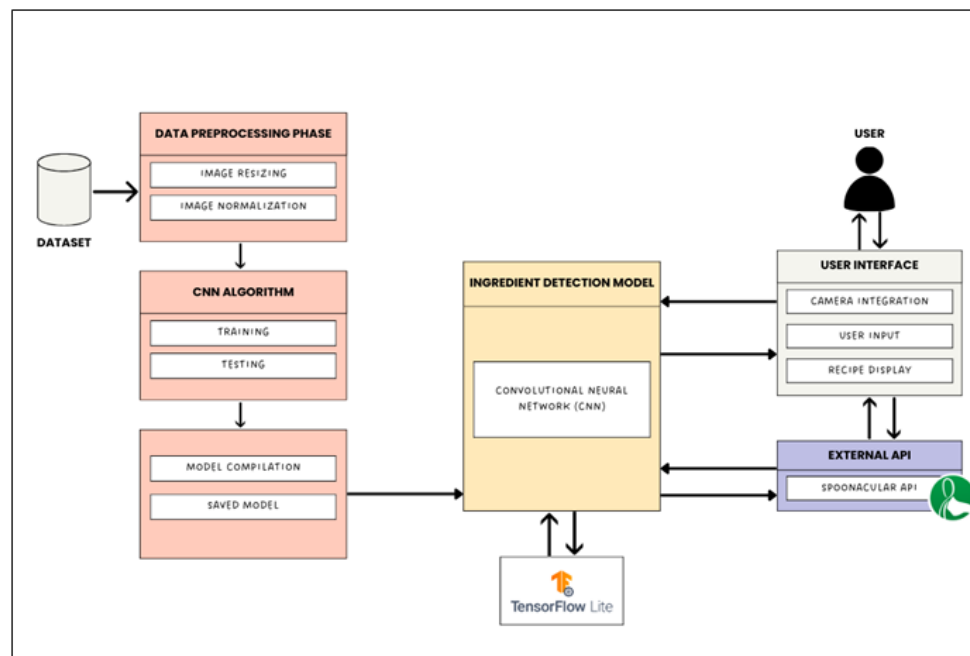


Fig. 3. System architecture of the project

3.5 Use case diagram

A use case diagram is constructed to identify the relationship between users and functionalities involved in the application. Fig. 4 shows the use case diagram designed for this application. The use case diagram outlines the interaction between the user and the application, highlighting the primary functions and processes involved. Based on Fig. 3, the process begins when the user Launches the Application. This initial step involves opening the app on their mobile device, which brings up the main interface of the food recipe assistant. From here, the user has several options for interacting with the system. The user can choose to search recipes manually, Scan Ingredients using Camera, or Upload Image. The search recipe manually option allows users to type in keywords or browse through a list of recipes based on their preferences. This is useful for users who already have a specific recipe in mind or want to explore new recipes. Alternatively, the Scan Ingredient using Camera option enables users to take a photo of an ingredient directly within the app. This image is then processed by the application to identify the ingredient using Convolutional Neural Network (CNN) technology. The upload image function allows users to select a pre-existing photo of an ingredient from their device's gallery. This flexibility ensures that users can easily provide the necessary input data for ingredient recognition. Once an image is provided through either scanning or uploading, the system proceeds to recognize ingredients using CNN. The CNN model analyzes the image to accurately identify the ingredient, leveraging its trained dataset and recognition capabilities. If the ingredient is successfully recognized, the application will Provide Recipe Recommendation based on the identified ingredient. This step involves querying the database for recipes that include the recognized ingredient, ensuring that users receive relevant and practical recipe suggestions. Users can then view the recipe, which provides detailed instructions and ingredient lists for preparing the recommended dishes. In the event of any issues, such as unrecognized ingredients or technical errors, the system includes an Error Handling process to manage these situations and ensure a smooth user experience.

<https://doi.org/10.24191/jcrinn.v10i2.554>

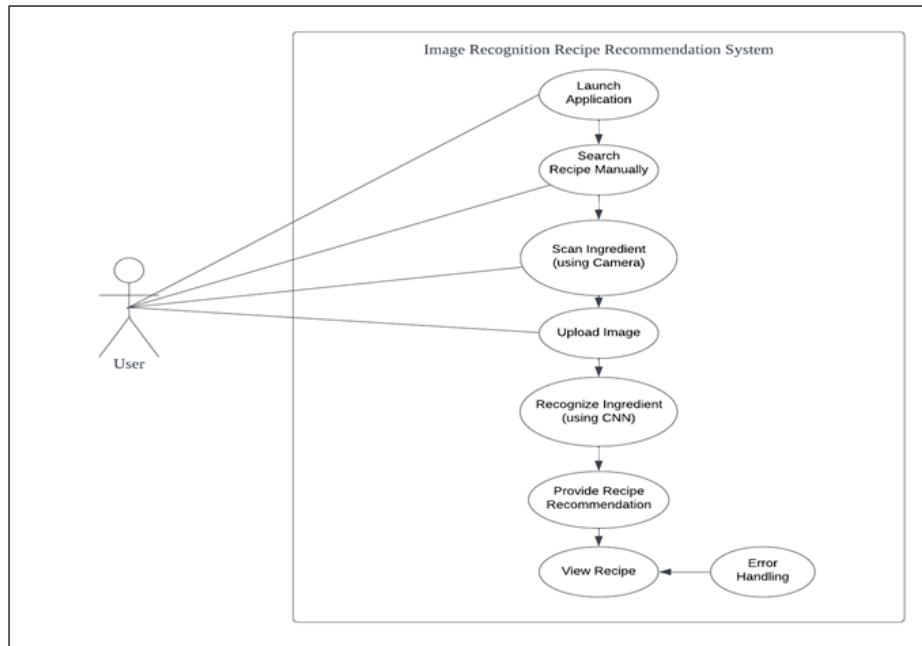


Fig. 4. Use case diagram for the application development

4. RESULTS AND DISCUSSION

This section covers the evaluation of CNN model using data split ratio, graphical analysis of the training and validation accuracy over epochs and the confusion matrix results. The user interfaces of the prototype are presented at the last part of the section.

4.1 Evaluation of CNN model

The evaluation of the CNN model involved testing its development using three different data splitting ratios to determine which configuration yielded the highest accuracy. The chosen splits were 70:30, 80:20, and 90:10, and each was tested over 30 epochs to maintain consistency. The goal was to identify the optimal ratio that would balance the training and validation data, ensuring that the model generalizes well to unseen data. Table 2 shows the data split ratio and the results.

Table 2. Data split ratio and results

Data Splitting Ratio	Test Accuracy	Model Accuracy
80:20	0.870	95.8%
90:10	0.917	96.8%
70:30	0.925	96.8%

Based on Table 2, these results indicate that the 70:30 split provided the highest test accuracy, closely followed by the 80:20 and 90:10 splits. However, the differences in accuracy between the splits were relatively small, suggesting that the model performs consistently well across different training-validation splits. This consistency is critical for ensuring the robustness of the CNN model in real-world applications. Fig. 5 shows the epochs for the testing using 70:30 split.

```

Epoch 1/30
260/260 [=====] - 108s 411ms/step - loss: 0.5266 - accuracy: 0.8094 - val_loss: 0.3334 - val_accuracy: 0.8619
Epoch 2/30
260/260 [=====] - 106s 410ms/step - loss: 0.2999 - accuracy: 0.8863 - val_loss: 0.2592 - val_accuracy: 0.8937
Epoch 3/30
260/260 [=====] - 107s 410ms/step - loss: 0.2176 - accuracy: 0.9208 - val_loss: 0.1980 - val_accuracy: 0.9246
Epoch 4/30
260/260 [=====] - 106s 409ms/step - loss: 0.1763 - accuracy: 0.9329 - val_loss: 0.1908 - val_accuracy: 0.9322
Epoch 5/30
260/260 [=====] - 106s 407ms/step - loss: 0.1379 - accuracy: 0.9485 - val_loss: 0.2214 - val_accuracy: 0.9151
Epoch 6/30
260/260 [=====] - 108s 415ms/step - loss: 0.1178 - accuracy: 0.9547 - val_loss: 0.1771 - val_accuracy: 0.9367
Epoch 7/30
260/260 [=====] - 108s 414ms/step - loss: 0.0972 - accuracy: 0.9656 - val_loss: 0.1714 - val_accuracy: 0.9432
Epoch 8/30
260/260 [=====] - 105s 405ms/step - loss: 0.0789 - accuracy: 0.9699 - val_loss: 0.1870 - val_accuracy: 0.9303
Epoch 9/30
260/260 [=====] - 105s 406ms/step - loss: 0.0744 - accuracy: 0.9742 - val_loss: 0.1917 - val_accuracy: 0.9491
Epoch 10/30
260/260 [=====] - 106s 409ms/step - loss: 0.0614 - accuracy: 0.9776 - val_loss: 0.1955 - val_accuracy: 0.9483
Epoch 11/30
260/260 [=====] - 106s 408ms/step - loss: 0.0420 - accuracy: 0.9846 - val_loss: 0.2151 - val_accuracy: 0.9494
Epoch 12/30
260/260 [=====] - 113s 436ms/step - loss: 0.0335 - accuracy: 0.9878 - val_loss: 0.4647 - val_accuracy: 0.9125
Epoch 13/30
...
Epoch 17/30
260/260 [=====] - 104s 402ms/step - loss: 0.0261 - accuracy: 0.9918 - val_loss: 0.3179 - val_accuracy: 0.9317
25/25 [=====] - 4s 164ms/step - loss: 0.2743 - accuracy: 0.9250
Test accuracy for 30-70 split: 0.925000011920929
Output is truncated. View as a scrollable element or open in a text editor. Adjust cell output settings--

```

Fig. 5. Epochs for 70:30 split testing

Early stopping was used in the model-building process to prevent overfitting and guarantee effective training. Early stopping keeps track of how well the model is performing on the validation set and stops training when it no longer improves. This method helps prevent overfitting, which occurs when the model learns the noise in the training data rather than the underlying patterns and is especially helpful when training for a large number of epochs. An extra measure of protection was offered by the early stopping method, which made sure that the model training would end if more epochs did not result in performance improvements. This prevented overfitting and saved computing resources.

4.2 Analysis of results and graphical representation

Graphical analysis of the training and validation accuracy over epochs for each data split reveals valuable insights into the model's learning process. The training accuracy curves consistently approached high values, indicating that the model learned effectively from the training data. The validation accuracy curves, meanwhile, demonstrated the model's ability to generalize to unseen data. Fig. 6 shows the graph of 70:30 model split, which has shown the best results.

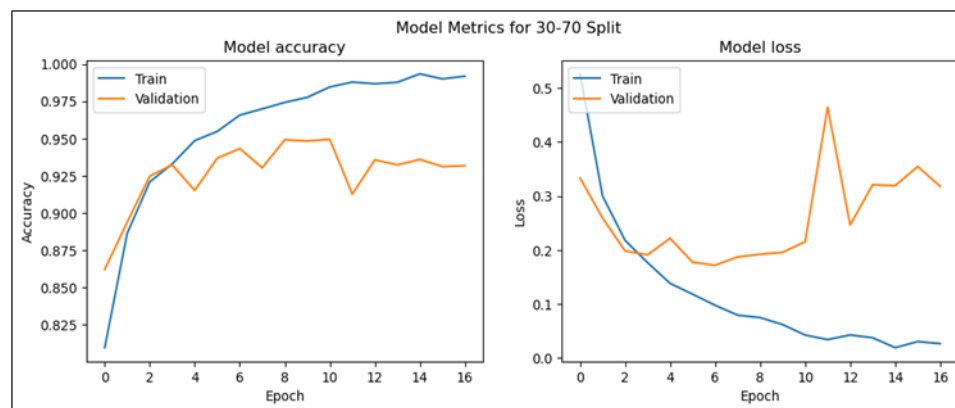


Fig. 6. Graph of 70:30 model split

The graphs in Fig. 6 display the model metrics for a 70-30 training-validation split, showing trends in accuracy and loss over 16 epochs. In the accuracy graph, the training accuracy (blue line) steadily increases, approaching near-perfect accuracy as epochs progress, indicating that the model is learning effectively from the training data. The validation accuracy (orange line), while generally improving, shows some fluctuations, which may suggest slight overfitting as the model starts to memorize the training data rather than generalizing well. In the loss graph, the training loss (blue line) consistently decreases, reflecting the model's improving performance. The validation loss (orange line), however, shows a more variable trend with occasional spikes, further indicating potential overfitting issues. Despite these fluctuations, the overall trend suggests that the model is learning and improving, but adjustments in hyperparameters or additional regularization may be needed to enhance generalization to new data.

Fig. 7 shows the graphs which display the model metrics for a 80-20 training-validation split, showing trends in accuracy and loss over 12 epochs. In the accuracy graph, the training accuracy (blue line) shows a steady increase, indicating effective learning from the training data, and reaches high levels by the end of the epochs. The validation accuracy (orange line), while initially improving, shows some fluctuations, suggesting that the model is experiencing slight overfitting as it starts to memorize the training data rather than generalizing. In the loss graph, the training loss (blue line) consistently decreases, reflecting the model's improving performance. The validation loss (orange line) initially decreases but then fluctuates, indicating variability in how well the model generalizes to unseen data. Despite these fluctuations, the overall trend suggests that the model is learning well but may benefit from adjustments in hyperparameters or regularization techniques to enhance its generalization capabilities and reduce overfitting.

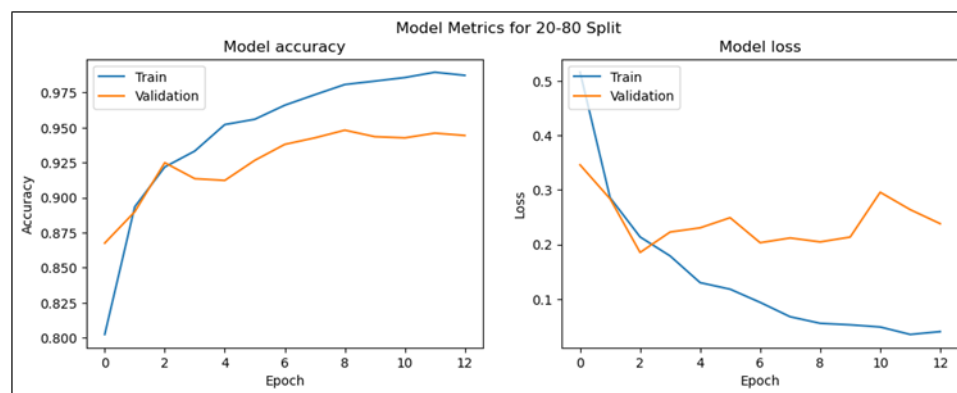


Fig. 7. Graph of 80:20 model split

Fig. 8 shows the graphs which display the model metrics for a 90-10 training-validation split, showing trends in accuracy and loss over 12 epochs. In the accuracy graph, the training accuracy (blue line) steadily increases, indicating effective learning from the training data and approaching near-perfect accuracy. The validation accuracy (orange line) also shows an overall upward trend, although it exhibits some fluctuations, suggesting minor overfitting but still maintaining high performance. In the loss graph, the training loss (blue line) consistently decreases, reflecting the model's improving ability to minimize errors during training. The validation loss (orange line), while initially decreasing, shows some variability, indicating fluctuations in the model's generalization to unseen data. Despite these fluctuations, both accuracy and loss trends suggest that the model is performing well, but further fine-tuning of hyperparameters or regularization techniques may be beneficial to enhance its stability and generalization.

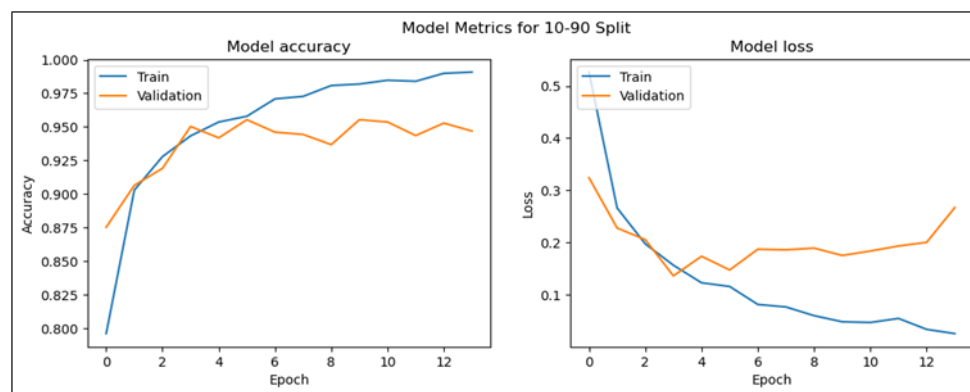


Fig. 8. Graph of 90:10 model split

Based on the overall results, the 70:30 split, the model showed a slight edge in test accuracy, suggesting that having more validation data helps in fine-tuning the model's hyperparameters and ensuring better generalization. The 80:20 split provided a balanced approach, yielding high test accuracy while maintaining substantial training data. The 90:10 split, while slightly less effective in test accuracy, still produced robust results, highlighting the model's capacity to learn well even with a smaller validation set.

In conclusion, the evaluation of the CNN model through different data splitting ratios and the implementation of early stopping have proven to be effective strategies for developing a robust ingredient recognition system. Due to the early stopping, the number of epochs conducted has been different in the data split ratio testing. The number of epochs in machine learning training differs when using early stopping because early stopping is a dynamic, adaptive mechanism designed to prevent overfitting and unnecessary training. In this research, the 70:30 split emerged as the most optimal configuration, providing the highest test accuracy. However, the consistent performance across different splits underscores the reliability of the CNN model in accurately identifying ingredients, paving the way for its successful deployment in the food recipe assistance application.

4.3 Confusion matrix results

The accuracy testing for the ingredient detection system involves the utilization of a confusion matrix reflecting the model's classification performance on 5 datasets with over 14,000 images. The matrix comprises four key metrics: True Positive (TP), False Positive (FP), False Negative (FN), and True Negative (TN). TP represents instances where the model correctly identifies ingredients within the 5 classes, FP denotes instances where the model incorrectly identifies non-ingredients as ingredients, FN indicates instances where the model incorrectly identifies ingredients as non-ingredients, and TN signifies instances where the model correctly identifies non-ingredients. Fig. 9 show the results of the confusion matrix for 70:30 data split.

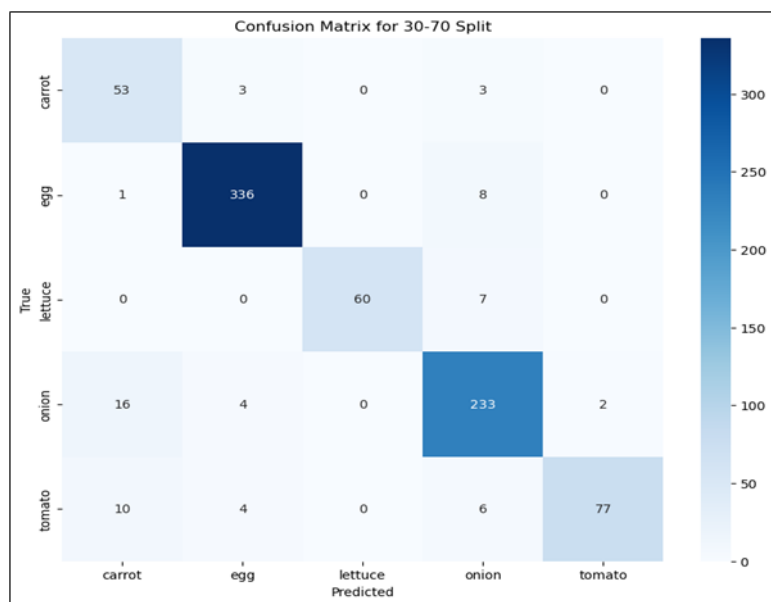


Fig. 9. Confusion matrix results

For the specific CNN model being evaluated, the confusion matrix in Fig. 9 shows the following values for different classes. The model correctly identified 53 instances of carrot but misclassified 3 as egg, 3 as onion, and 0 as tomato. For the egg class, the model accurately predicted 336 instances, while 1 was misclassified as carrot, 0 as lettuce, 8 as onion, and 0 as tomato. Lettuce was correctly identified in 60 instances, with 0 misclassified as carrot, egg, and 7 as onion. Onion had 233 correct predictions, but 16 were misclassified as carrot, 4 as egg, and 2 as tomato. Lastly, the tomato class had 77 correct predictions, with 10 misclassified as carrot, 4 as egg, and 6 as onion.

Based on these values, the accuracy, precision, recall, and F1 score are calculated to evaluate the performance of the CNN model. The overall accuracy of the CNN model is determined to be 0.92, indicating that 92% of the model's predictions are correct. For the egg class, precision is calculated as 96.6%, indicating that 96.6% of the samples identified as egg by the model were indeed egg. The recall for the egg class is also 96.6%, showing that the model accurately identifies egg 96.6% of the time when predicting a sample as egg. The F1 score for the egg class is 96.6%, providing a balanced measure of the model's performance for this class. These metrics collectively provide a comprehensive evaluation of the model's performance across different classes. Table 3 shows the results of the confusion matrix results.

Table 3. Calculation of accuracy, precision, recall, F1 score

	Calculation	Answer	Percentage
Accuracy	$(TP+TN) / (TP+TN+FP+FN)$ $(152+452) / (152+452+193+193)$	0.968	96.8%
Precision	$TP / (TP+FP)$ $152 / (152+193)$	0.842	84.2%
Recall	$TP / (TP+FN)$ $152 / (152+193)$	0.970	97.0%
F1 Score	$2 * (precision * recall) / (precision + recall)$ $2(0.441 * 0.441) / (0.441 + 0.441)$	0.896	89.6%

The results of the system are consistent with the calculations done manually, as shown in Fig. 9. The system's calculations for metrics like accuracy, precision, recall, and F1 score are verified by the consistency of the results. The accuracy and dependability of the implemented ingredient detection system are highlighted by the match between the values generated by the system and the calculations done manually. This validation ensures the validity of the evaluation metrics used in the project and increases confidence in the system's ability to accurately assess the Convolutional Neural Network model's performance.

4.4 System usability scale (SUS) results

Fig. 10 shows the histogram representing the distribution of SUS scores for the mobile application. The SUS assessment was conducted to evaluate the usability of the mobile application. A total of 20 respondents participated in the evaluation and the SUS scores ranged from approximately 55 to 100, with the average score recorded at 83.33. This average score indicates that the application demonstrates excellent usability, with strong acceptance and satisfaction among users. The histogram reveals that the majority of users rated the system between 80 and 95, with the highest concentration of scores occurring in the 90–95 range. This suggests that most users found the system to be intuitive, efficient, and satisfying to use. Only one respondent provided a relatively lower score, which appears to be an outlier. This might indicate a unique usability issue or a different user expectation, but not a critical issue. Overall, the results suggest that the application is well-designed and user-friendly, with strong user satisfaction and minimal usability issues.

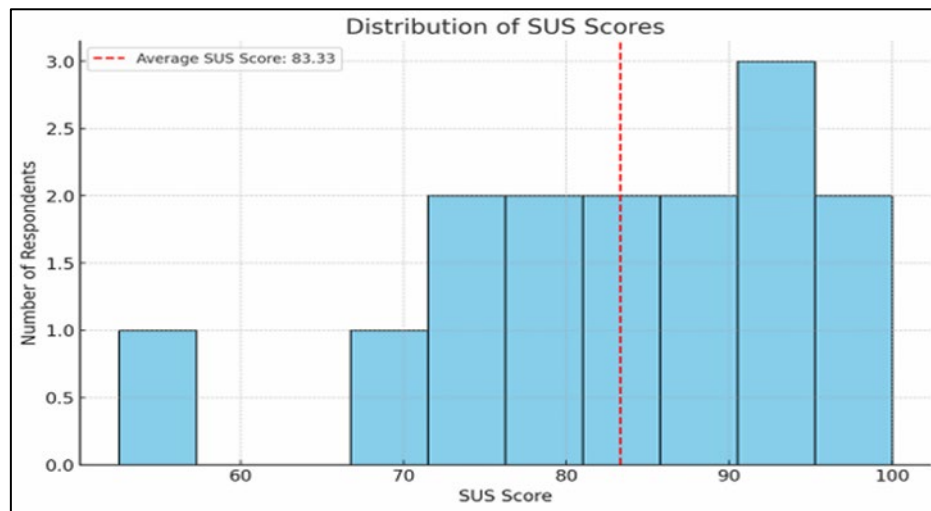


Fig. 10. SUS score results

4.5 User interface of the mobile application

Microsoft Visual Studio's Flutter is used to create the user interface for the "CNN Integrated Mobile Application for Food Recipe Assistance" prototype. The app's homepage's graphical user interface (GUI) is shown in Fig. 11. When the homepage loads, the majority of the white space is occupied by the camera screen. This bold design decision was made on purpose to draw attention to the app's core function, which is ingredient detection. Towards the bottom of the page are three primary buttons: an "image picker," a "camera icon," and a "search icon".

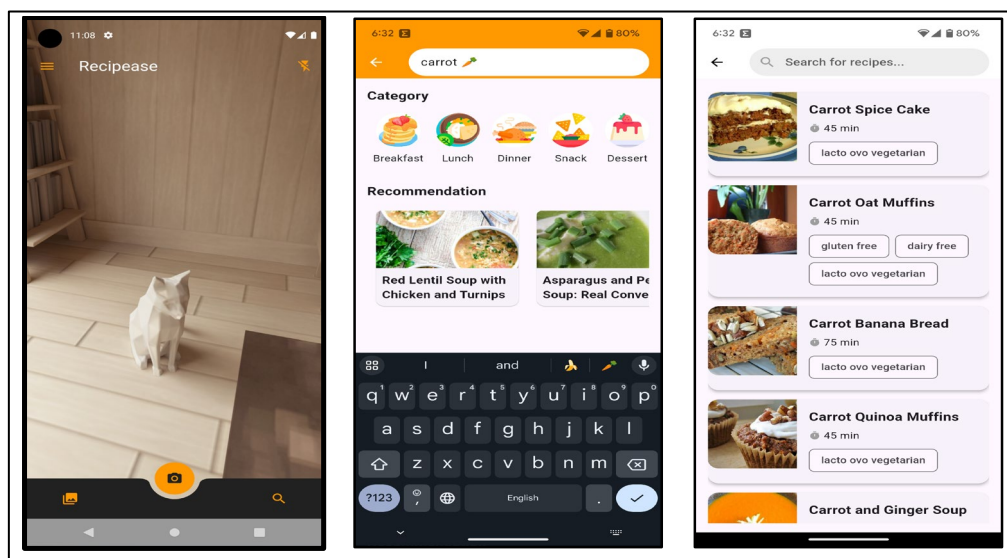


Fig. 11. Homepage and the search page of the application

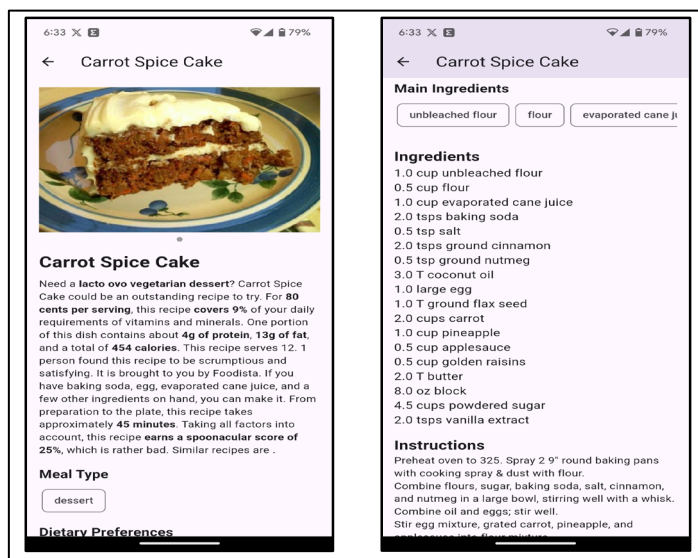


Fig. 12. Sample recipe based on menu

As shown also in Fig. 11, pressing the search icon brings up the search page and applies the Spoonacular API. Users can manually search for recipes on this page by entering keywords into the search bar. The keywords are not limited to name of recipe or ingredient used but even cuisine, dietary preferences and even meal types such as breakfast, lunch or dinner. This feature makes sure that even if consumers decide not to use the picture detection functionality, they can still find recipes. Fig.12 shows the sample of recipe details for Carrot Spice Cake when user clicks on the menu.

In the meantime, the CNN model for ingredient detection is directly connected to the camera icon and image selection button. The software examines the image and goes to the statistics page when an image is selected or a snapshot is taken. Then, it shows the detected ingredient and its confidence level. An example of the statistics page with "carrot" selected as the detected ingredient is displayed in Fig. 10. Users can view the ingredient that the CNN model identified in detail on this page, along with the confidence level and the image that was utilized for identification. As illustrated also in Fig. 13, if an ingredient or item is entered that does not fall into one of the five predefined classes, the system correctly detects the lack of a recognized ingredient and displays a pop-up message indicating a detection error. This visual aid for error handling makes sure users are aware of any problems with their input and improves the user experience overall by giving quick, clear feedback.

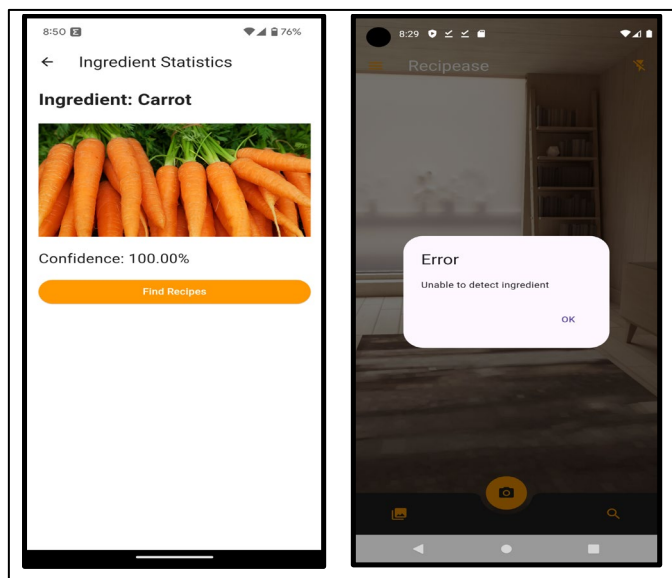


Fig. 13. Statistics page with carrot as input and pop-up message of error handling

5. CONCLUSION AND RECOMMENDATION

CNN Integrated Mobile Application for Food Recipe Assistance has effectively accomplished its goals by utilizing cutting-edge machine learning methods to improve component identification and offer customized recipe suggestions. The experiment showed how Convolutional Neural Networks (CNN) can be used effectively to reliably identify items from user-provided photos, which will help with meal preparation efficiency and cut down on food waste. The significance of the research is that it is in perfect harmony with the Sustainable Development Goal (SDG) 12 of the United Nations, which places a strong emphasis on responsible production and consumption. It promotes effective ingredient management, enabling consumers to maximise the amount of food in their pantries and drastically lowering the possibility of needless food waste. This improves everyday cooking's usefulness and supports the application's objective of encouraging responsible consumption. The software supports a more sustainable and thoughtful approach to home cooking by providing people with the tools to use ingredients effectively. The main contribution of the research project is the food ingredient image recognition using CNN that is implemented for mobile platform. Most of the current similar research project are using web based platform which has

its limitation in mobile application accessibility, such as poor screen reader support or low contrast and small text.

This mobile application was developed and implemented using an organized methodology, guaranteeing a reliable user experience. Thorough testing and assessment established the good accuracy of 96.8% and dependability of the application, fulfilling the project's main objectives. Nevertheless, the project encountered many constraints that offer prospects for forthcoming enhancements. This initial project only focus on 5 common ingredients, which limits its usefulness for users with a variety of culinary requirements. Furthermore, relying on outside APIs to provide recipe recommendations could lead to dependency and possible restrictions on the consistency and availability of data. Future improvements should focus on broadening the dataset to include a wider variety of ingredients, improving the CNN model's training to increase accuracy, and adding more features like meal planning tools and nutritional data. Furthermore, investigating the incorporation of offline features and diminishing reliance on external APIs may enhance the robustness of the program and enhance user contentment. These enhancements will not only address the current limitations but also pave the way for a more comprehensive and versatile solution, solidifying the project's contribution to the culinary and technology fields.

6. ACKNOWLEDGEMENTS/FUNDING

The authors would like to acknowledge the support of UiTM Cawangan Terengganu for the continuous support given to the research initiatives in the university.

7. CONFLICT OF INTEREST STATEMENT

The authors declare that this research was conducted in the absence of any conflict of interests.

8. AUTHORS' CONTRIBUTIONS

Muhammad Imran Nor Azlan Shah: Development and analysis, writing original draft; **Norlina Mohd Sabri:** Supervision and editing draft; **Gloria Jennis Tan:** Review manuscript; **Zhiping Zhang:** Formatting and final review.

9. REFERENCES

- Anand, L., Tyagi, R., & Mehta, V. (2024). Food recognition using deep learning for recipe and restaurant recommendation. In Bhateja, V., Lin, H., Simic, M., Attique Khan, M., Garg, H. (Eds) *Cyber Security and Intelligent Systems. Lecture Notes in Networks and Systems* (pp. 1056). Springer, Singapore. https://doi.org/10.1007/978-981-97-4892-1_23
- Anitha, E., Nandhini, S., Palani, H. K., & Marshal, M. C. (2023). Recipe recommendation using image classification. In *Proceedings of the 2023 5th International Conference on Smart City Applications* (pp. 1023–1033). Springer International Publishing.
- Boyd, L., Nnamoko, N., & Lopes, R. (2024). Fine-grained food image recognition: A study on optimising convolutional neural networks for improved performance. *Journal of Imaging*, 10(6), 126. <https://doi.org/10.3390/jimaging10060126>
- Chen, Y. C., & Chiang, H. C. (2025). Deep learning-based automatic food identification with numeric label. *Multimedia Tools and Applications*. <https://doi.org/10.1007/s11042-025-20648-x>

- Chopra, B., Jain, G., Mahendra, N., Sinha, S., & Natarajan, S. (2023). Ingredient detection, title and recipe retrieval from food images. In *International Conference on Ubiquitous and Future Networks, ICUFN* (pp. 288-293). IEEE. <https://doi.org/10.1109/ICUFN57995.2023.10199735>
- Dsouza, H. J., Nayak, K. S., Karkera, K. K., Varghese, M., & Shreejith, K. B. (2024). Ingredient detection and recipe recommendation using deep learning. *International Journal of Advanced Research in Computer and Communication Engineering*, 13(3), 630–635. <https://ijarccce.com/wp-content/uploads/2024/04/IJARCCCE.2024.133100.pdf>
- Gautam, G., & Khanna, A. (2024). Content based image retrieval system using cnn based deep learning models. *Procedia Computer Science*, 235(2023), 3131–3141. <https://doi.org/10.1016/j.procs.2024.04.296>
- Kansaksiri, P., Panomkhet, P., & Tantisuwichwong, N. (2023). Smart cuisine: Generative recipe & ChatGPT powered nutrition assistance for sustainable cooking. *Procedia Computer Science*, 225, 2028–2036. <https://doi.org/10.1016/j.procs.2023.10.193>
- Kurian, V., & Jacob, V. (2023). Importance of CNN in the classification of remote sensing images. In *Proceedings of the ACCTHPA 2023 - Conference on Advanced Computing and Communication Technologies for High Performance Applications* (pp. 1-4). IEEE. <https://doi.org/10.1109/ACCTHPA57160.2023.10083375>
- Li, Z. B. (2023). Predict food recipe from image using CNN/transformer & knowledge distillation for portability. In *2023 IEEE International Conference on Sensors, Electronics and Computer Engineering* (pp. 852–859). IEEE. <https://doi.org/10.1109/ICSECE58870.2023.10263583>
- Navaprakash, N., Reddy, S. V., & Dakshinesh, S. (2025). Improving accuracy in text extraction from images using region-based convolutional neural networks algorithm compared to convolutional neural network algorithm. In *2025 International Conference on Artificial Intelligence and Data Engineering (AIDE)* (pp. 706-710). IEEE. <https://doi.org/10.1109/AIDE64228.2025.10987308>
- Rodrigues, M. S., Fidalgo, F., & Oliveira, Â. (2023). RecipeIS—Recipe recommendation system based on recognition of food ingredients. *Applied Sciences (Switzerland)*, 13(13), 1-19. <https://doi.org/10.3390/app13137880>
- Singh, P. K., & Susan, S. (2023). Transfer learning using very deep pre-trained models for food image classification. In *2023 14th International Conference on Computing Communication and Networking Technologies* (pp. 1-6). IEEE. <https://doi.org/10.1109/ICCCNT56998.2023.10307479>
- Sri, B. R., & Balakrishnan, T. S. (2024). *Classification of pests in agricultural farms using convolutional neural network compared to artificial neural network*. 2024 15th International Conference on Computing Communication and Networking Technologies (ICCCNT), 1–4. <https://doi.org/10.1109/ICCCNT61001.2024.10725825>
- Swain, M., Manyatha, A. R., Dinesh, A. S., Sampatrao, G. S., & Soni, M. (2023). Ingredients to recipe: A YOLO-based object detector and recommendation system via clustering approach. In *Proceedings of the 3rd International Conference on Artificial Intelligence and Smart Energy* (1397-1404). IEEE. <https://doi.org/10.1109/ICAIS56108.2023.10073769>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Enhancing Pineapple Cultivar Classification: A Framework for Image Quality, Feature Extraction, and Algorithmic Refinement

Mohamad Faizal Ab Jabal¹, Muhammad Irfan Rozlan², Azrina Suhaimi^{3*},
Harshida Hasmy⁴

^{1,3,4}*Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Johor Branch, Pasir Gudang Campus, 81750 Masai, Johor, Malaysia.*

²*Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Terengganu Branch, Kuala Terengganu Campus, 21080, Kuala Terengganu, Terengganu, Malaysia.*

¹*Applied Mathematics & System Development (AMSys)-Special Interest Group (SIG), Universiti Teknologi MARA Johor Branch, Pasir Gudang Campus, 81750 Masai, Johor, Malaysia.*

ARTICLE INFO

Article history:

Received 22 June 2025

Revised 14 August 2025

Accepted 14 August 2025

Published 1 September 2025

Keywords:

Autonomous Agriculture

Colour Feature

Feature Extraction

Pineapple Classification

Pineapple Cultivar

DOI:

10.24191/jcrinn.v10i2

ABSTRACT

Accurate classification of pineapple cultivars is hindered by limitations in image acquisition, feature extraction, and classification algorithms. This study identifies these technical challenges and proposes methodological improvements to support reliable autonomous recognition systems. Key findings reveal deficiencies in current image datasets, leading to a proposed standardised acquisition protocol involving consistent lighting, optimal camera positioning, and suitable file formats. The HSV colour space is validated as more effective for extracting skin features, with its threshold values and post-processing steps significantly reducing classification errors. The proposed algorithmic refinements integrate chromatic and morphological attributes, particularly surface area and optimise logical operators to enhance accuracy. The study addresses two main research gaps: precise quantification of skin colour and the development of robust classification frameworks. Future work will emphasise empirical validation and the deployment of the YOLOv7 model for real-time, on-site assessment of fruit maturity in field conditions. These contributions hold strong implications for advancing precision agriculture and improving post-harvest processing.

1. INTRODUCTION

One of the main economic drivers in many nations is the agricultural sector, and pineapple is one of the most produced and traded tropical fruits in the world (Chang et al., 2025). Despite the crop's economic importance, most pineapple cultivars are still classified manually (Lai et al., 2023). Accurately identifying

^{3*} Corresponding author. *E-mail address:* azrin253@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.546>

pineapple cultivars is essential for maintaining efficient supply chain operations and meeting consumer demand. Harvesting and sorting pineapples continue to be labour-intensive processes, putting significant strain on the industry's manpower and expertise. Reliance on conventional farming practices lowers operational efficiency, which is made worse by labour shortages brought on by rural migration (Lai et al., 2023; Li et al., 2024). These limitations impede producers' capacity to fulfil increasing market demand and limit the overall total growth potential of the pineapple industry (Lai et al., 2023; Li et al., 2024).

Existing research has predominantly focused on detecting ripeness rather than identifying specific pineapple cultivars. In this regard, Lai et al. (2023) created a system employing You Only Look Once version 7 algorithm (YOLOv7) to detect pineapple maturity under difficult conditions, whereas Trinh and Nguyen (2023) used YOLOv5 for real-time identification with high accuracy suitable for embedded devices to detect pineapple freshness. Additionally, Arboleda et al. (2021) used a fuzzy logic and image processing-based system to automatically assess pineapple ripeness of the 'queen' kind, increasing evaluation efficiency. Though maturity categorisation may aid in harvesting at the optimal time for optimum market value, it does not address the problem of cultivar identification, which is related with a distinct technique based on visual traits. Moreover, Logistic Regression (LR), Naïve Bayes (NB), and Support Vector Machine (SVM) have been successful in image classification applications (Wang & Chai, 2022). However, using handcrafted features during extraction poses significant hurdles for these models, making them less adaptive to complicated agricultural data.

The goal of this research is to identify a reliable approach for classifying pineapple cultivars. Along with the discussion in this paper, the challenges, advantages, and disadvantages of existing approaches will be emphasized. Furthermore, this paper proposes potential improvements to current classification approaches. The significance of this research is to highlight the research gap and opportunities for development in the classification domain. Furthermore, the findings will aid other researchers in developing an automatic pineapple cultivar detection and classification system, so that predictions about the food supply chain can be visualized and the food supply is secure. Indirectly, secure the nation's economy, which is dependent on pineapple imports and exports.

In accordance with the study's objectives, this paper is structured as follows: Section II presents a comprehensive literature review and critical analysis of related work. Section III details the proposed methodology. Section IV examines the analytical findings and provides an in-depth discussion of the results. Finally, Section V concludes the study and proposes directions for future research.

2. LITERATURE REVIEW

2.1 Pineapple feature

Features are essential data used to recognize and classify objects in images. The features accessible for object recognition and classification include shape, size, texture, and colour (Jabal et al., 2022, Jabal et al., 2024). Consequently, analysis of pineapple images suggests that shape and colour are the most reliable features for use in this study. Table 1 displays the findings from the analysis.

Table 1. Pineapple features

Cultivar	Features	Other Pineapple Name called by Local People's
Spanish	1. Pointed leaf ends with spines. 2. Produces 2-12 basal slips at the stalk base. 3. Ideal for canning. 4. Ripe fruit changes from green to dark purple or reddish-orange. 5. Large crown. 6. Long shelf life. 7. Cylindrical fruit shape.	1. Singapore Spanish Pineapple. 2. Selangor Green Pineapple. 3. Gandul Pineapple.

Smooth Cayenne	1. Large fruit size. 2. Flat pineapple eyes. 3. Leaves are dark green. 4. Tapered fruit shape.	1. Sarawak Pineapple.
Queen	1. Tapered fruit shape. 2. Best consumed fresh. 3. Leaves are bluish-green with purple centres. 4. Spiny leaves.	1. Moris Pineapple. 2. Yankee Pineapple. 3. Moris Elephant Pineapple.
Hybrid	1. Cylindrical fruit shape. 2. Flat pineapple eyes. 3. Green leaves. 4. Slightly spiny leaves. 5. Best consumed fresh.	1. Maspine Pineapple. 2. Josapine Pineapple. 3. N36 Pineapple. 4. MD2 Pineapple.

Based on Table 1, features related to the crown size, leaves and the practicality of the pineapple in accommodation with the post-harvest handling come into play as major aspects contributing to the marketing and production viability of the cultivar. Various cultivars possess certain features that lend themselves to uses such as canning, fresh use, fresh-cut products. Information on agronomics is to optimising production, transport, and trade. Furthermore, Fig. 1 shows the sample of the pineapple images received from the Malaysian Pineapple Industry Board (2024). Based on Fig. 1, it contains eight sample images of pineapple according to their cultivar. The shape and colour are evident qualities that can be utilized to recognise and classify. Furthermore, the size of the pineapple might be considered as an additional feature to guarantee that it is classified accurately.

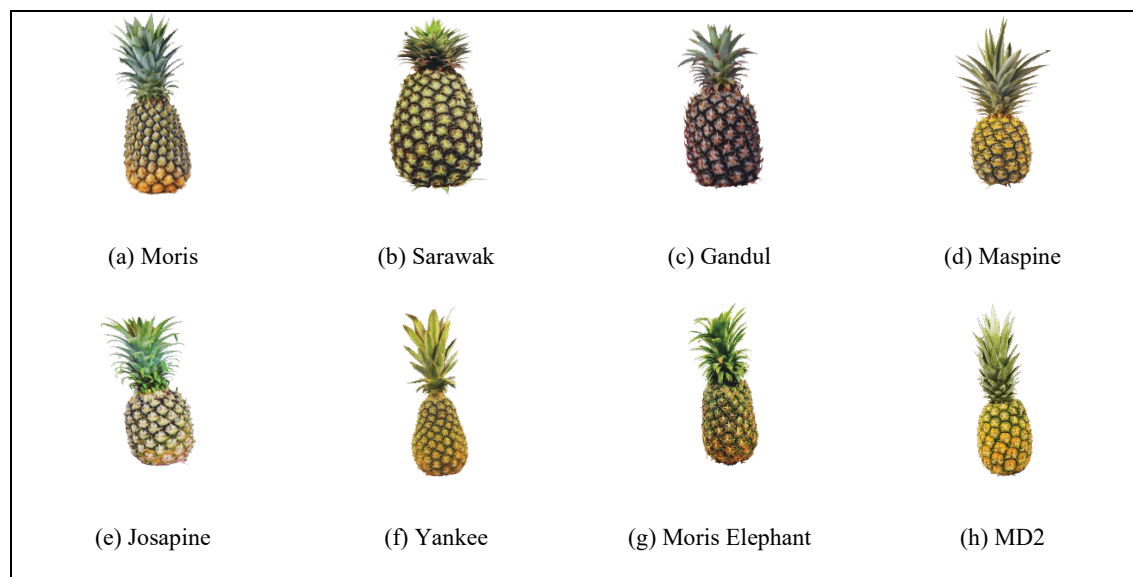


Fig. 1. Pineapple cultivar images

Next, this study discusses potential recognition approaches proposed by other researchers. The extensive discussion will emphasize the processes involved, advantages and disadvantages of the approaches.

2.2 Feature extraction approach

Feature extraction is the crucial process before the object in the image can be classified (Jabal et al., 2024). Hence, the accuracy of the object classification result is influenced by the extraction result. Several approaches have recently been developed to extract significant features such as colour, shape, and texture that aid in the classification of pineapple cultivars or maturity levels. Such processes determine both the accuracy of the method and its practical application in real-world agricultural contexts.

Lai et al., (2023) has deployed the integration of YOLOv7 and SimAM to extract the shape and colour of the object. The colour space used to extract the colour feature is Red, Green and Blue (RGB). The authors have done a modification on MConv structure and soft-NMS algorithm to extract the accurate result. However, the achievement of the accuracy rate is not recorded. Furthermore, the proposed approach had required high performance workstation and consumed high memory during execution. On the other hand, Manik et al., (2024) has used morphological method to extract 13 pineapple features. The first step is to eliminate the non-ROI area in the image. Second step the process continues to resize the image and segment the features according to the class. However, there is no record regarding the extraction accuracy rate reported by the authors.

In different cases, with the goal of the research to extract the shape, colour and texture feature, Wan Nurazwin et al., (2022) has implemented Analysis of Variance (ANOVA) to select and extract the crucial features from an image. However, the author does not report the achievement of the accuracy rate. Moreover, the approach needs to be improved to deal with the uncondusive foreground and colour similarities in the image, which caused the approach to perform incorrect feature extraction. Next, in 2021 Arboleda et al. has used RGB colour extraction to yield the colour feature. The author has reported that the extraction accuracy rate was 90%. However, the approach has an issue to deal with the uncondusive environment image.

In such cases, Jabal et al., (2024) has implemented the Hue, Saturation and Value (HSV) Histogram method to extract the colour feature. HSV colour space is used to deal with the complex foreground in the image. The extraction accuracy rate is 77.86% was recorded by the author. However, the method needs improvement to deal with the image background which tends to dark brown and black colour.

Based on the extensive discussion, the advantages and disadvantages of the existing feature extraction approaches have been highlighted. According to the report there is an improvement opportunity towards the feature extraction approach to extract the shape and colour features of the pineapple. The key point is the correct shape and colour value will influence the approach to extract the correct ROI features.

2.3 Classification approach

Proceeding with discussion, this section will extensively discuss the existing classification approaches. This discussion aims to identify the most reliable approach for pineapple classification.

Environmental considerations significantly impact the effectiveness of pineapple classification systems, especially in real-world agricultural applications. Lighting conditions, occlusion by leaves or other objects, and fruit-background similarity can all have a substantial impact on image quality and classification accuracy (Chen et al., 2023, Lai et al., 2023). Fig. 2 presents the sample pineapple image with an uncondusive environment. The images have been achieved from the Facebook *Pembekal Benih Nanas MD2* (KSF Enterprise MD2, 2024). These issues stem from the intrinsic complexity of orchard ecosystems and a lack of predictability, which may affect the required robust procedures for developing design guidelines.



Fig. 2. Unconductive environment pineapple image

Fig. 2 portrays the sunlight, leaves, shadows, the cover on top of the pineapple, and a pile of pineapples in the container all affect classification accuracy. Chen et al. (2023) and Mukhiddinov et al. (2022) agreed that differences in texture and colour can cause visual inconsistencies and hinder pineapple recognition. Furthermore, lighting variances in the dataset present significant issues, as addressed by Lai et al. (2023), who noted that variable lighting causes brightness variations, which may hide essential elements such as shape and colour, which are critical for classification.

Therefore, to deal with the unconductive environment image, Shiu et al., (2023) utilized Faster Region-based Convolutional Neural Network (R-CNN) and Mask R-CNN to detect plastic-covered pineapples. Although the technique achieved 97.46% classification accuracy rate and 73.9% of average precision rate (AP), detection performance was poor and necessitated high quality imagery from unmanned aerial vehicle (UAVs), restricting its use in resource-constrained contexts. On the other hand, Roslan et al., (2023) applied the same approach by Shiu, but for orange and apple classification. Roslan's image depicts the orchard environments. The study has successfully overcome problems caused by fluctuating lighting conditions and obstructed fruits. Moreover, the study successfully achieved 91.44% of AP. However, no classification accuracy rate has been reported.

In 2023, Trinh and Nguyen deployed an improvement version of YOLOv5 to classify ripe pineapple and achieved a 98% classification accuracy rate. Moreover, the approach had achieved more than 96% of mean average precision (mAP) rate. Aside from that, the approach required a large dataset for training to assure accurate classification. Although the approach detection speed is 9.2 milliseconds (ms) per image, it took a long time to train. In such circumstances, Lai et al., (2023) were successful in classifying pineapples into 3 maturity classes by introducing an enhancement to YOLOv7. Furthermore, the approach has used the unconductive environment pineapple images. The approach was successful achieved 95.82% mAP. However, no record can be discovered of the classification's accuracy rate.

In such cases, Mukhiddinov et al., (2022) proposed a deep CNN model and an optimized YOLOv4 to classify fruits and vegetables based on unconductive environment image. Furthermore, the study observed 73.5% and 72.6% AP for fruits and vegetables respectively based on 20 classes. However, there is no record found on the classification accuracy rate. Besides, in 2021, Arboleda et al. used Fuzzy Logic to classify pineapples into 4 classes: 1) unripe, 2) overripe, 3) underripe, and 4) ripe (mature). According to the study, the classification accuracy rate was 100% for classes 1 and 2, and 90% for classes 3 and 4. However, only the conductive environment images are required by the approach. This constraint stresses the opportunity to improve the approach in dealing with an unconductive environmental image.

Meanwhile in 2021, utilised an aerial image of the pineapple orchard captured by a UAV with main objective to detect and count the pineapple crown. Wan Nurazwin et al., (2021) had conduct a study on machine learning classifiers, which are Artificial Neural Network (ANN), Support Vector Machine (SVM), Random Forest (RF), Naïve Bayes (NB), Decision Tree (DT) and k-Nearest Neighbours (kNN) to investigate the reliable classification approach in order to classify between pineapple (fruit) and pineapple crown (non-fruit). Later, Variable Learning Rate Backpropagation algorithm (GDX) is one of the models in ANN known as ANN-GDX is used to detect and count the pineapple crown. The result yield from the study is exceeded above 90% of classification accuracy rate to identified crown. However, finding from the study emphasized that the approach had potential for misclassification due to the noise in the image and colour similarities between leaves and crowns. Besides, the approach requires manual annotation of training data. Additionally, to implement the approach requires high spatial-resolution images in order to generate accurate results. On the other hand, Cai et al., (2020) had proposed enhanced version of Single Shot Multi-Box Detector (SSD) to deal with unconducive environment image. The study's goal is to classify fruits into specific classes. Furthermore, the study successfully produced 92.4% of mAP. However, the classification accuracy rate unable to be found in the report. Likewise, the approach has a shortcoming in that it requires manual image annotation.

This section might be summarised by emphasizing the findings following a thorough discussion. The unconducive environment image remains as the main challenge to recognise and classify the pineapple. In terms of feature extraction, shadow removal on the pineapple has the potential to remove or change the colour feature on the pineapple skin. Furthermore, the identical hue of the leaves and a portion of the pineapple skin has the potential to induce inaccurate feature extraction and ROI. Other than that, using RGB colour space in the histogram method for noise reduction is less reliable than using HSV colour space. Furthermore, miscalculation on extracted feature data by the developed equations in each classification approach developed by the previous researchers remains a major challenge in achieving an accurate pineapple cultivar classification result. Also, require manual annotation and large numbers of image for training emphasizes the limitations of the approaches. Additionally, some of the approaches demand for the high specification of the workstation for execution.

On the other hand, this study had difficulty determining from prior studies the true properties of colour features of the pineapple skin to be employed for extraction. There are numerous variables for colour features that can affect extraction. According to Jabal et al., (2022), RGB colour space, where each colour is represented by an integer between 0 and 255. Moreover, various values will be produced by combining the RGB colour space. Hence, this number can cause miscalculation during the extraction and affect the classification result. Additionally, this issue had stressed the RGB colour space is not reliable to be used for classification purposes. Another element influencing the colour of the pineapple skin in the picture is the light illumination. Compared to photography studio illumination, taking a picture in the sun will provide a distinct colour quality. This issue has the effect of causing the various colour properties to be extracted from the image and inputted to the equation used to determine classification. Consequently, the classification output that is generated will be inaccurate. This issue was discovered during the image capturing, and it can be resolved by following the correct standard operation procedure (SOP).

Following a thorough discussion, this study moves on to the next section to perform further analysis on the highlighted issue.

3. METHODOLOGY

The unconducive environment image, undefined properties of the pineapple skin colour and standard operating procedure (SOP) are the issues concluded in the previous section. This section provides a further analysis of the highlighted issues. The discussion will begin with the analysis on the dataset and continued with the colour features. Finally, the discussion will be ended with the summary of the section.

3.1 Dataset

The dataset used in this study was obtained from the Malaysian Pineapple Industry Board (2024) and Facebook *Pembekal Benih Nanas MD2* (KSF Enterprise MD2, 2024) as presented in Fig. 1 and 2 respectively. Furthermore, for the analysis purposes, another dataset was achieved from Kaggle as depicted in Fig. 3 (Sujitra, 2024).

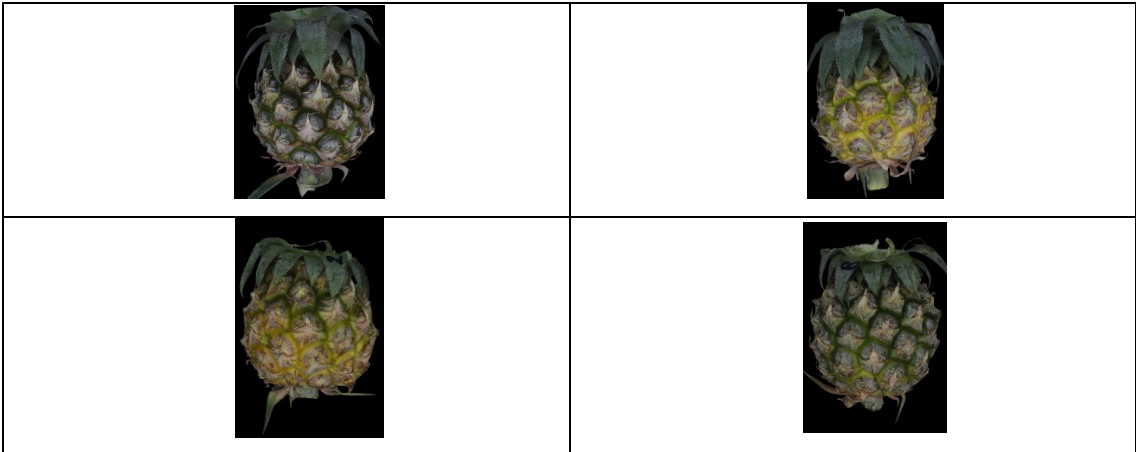


Fig. 3. Normal pineapple image

Source: Sujitra (2024)

Fig. 3 shows the sample pineapple images from Kaggle, which are the images were captured in the conducive environment same as shown in Fig. 1. However, the images in Fig. 2 contain noise, which may affect classification accuracy. Furthermore, it has been determined that the noises in the images are caused by the cover on the pineapple crown, pineapple leaves, and container used to collect the pineapple. Based on the analysis, the cover on the pineapple crown has caused the shadow on the pineapple skin. Besides, the leaves also created the shadow on the pineapple skin. Moreover, the similarity of the leaves colour might disrupt the colour of the pineapple skin that caused confusion during features extraction. Additionally, the container colour also caused disruptions during feature extraction. A pile of pineapples in the container is another factor that confuses the feature extraction, resulting in an inaccurate pineapple classification result. Next, the conducted analysis has examined the properties of the images as shown in Table 2.

Table 2. Sample of image properties

Image	Description
 (a)	Dimension: 295 × 638 pixels Bit Depth: 32 Format: PNG Colour: sRGB



(b)

Dimension: 720 × 960 pixels
 Bit Depth: 24
 Format: JPG
 Colour: sRGB



(c)

Dimension: 417 × 598 pixels
 Bit Depth: 32
 Format: PNG
 Colour: sRGB

Table 2 shows that the collection of datasets used in this study has a variety of image dimension and 2 types of file formats. The most common file format are Portable Network Graphic (PNG) and Joint Photographic Expert Group (JPG). Moreover, the colour for all images is standard Red, Green, and Blue (sRGB). Based on Table 2, the bit depth of the image is a key to determine the value of the RGB used for extraction. The higher of the bit depth, the higher the tones of the colour, indicating that the colour generated by the camera is similar to the actual colour visible to the human eye. Furthermore, the significance of the higher tones of the colour will increase the contradiction of distinguishing ROI features from non-ROI features in the image, even if the colours are slightly similar, such as the green colour of the leaves and the green colour of the pineapple skin. As a result, the 24-bit depth in image (b) may hinder ROI differentiation due to noise interference during feature extraction.

3.2 Colour features

The finding from the analysis towards the datasets, green and yellow are the dominant colours on the pineapple skin. However, the gradient of the green and yellow has difficulty the extraction when the various values have been generated from the RGB colour space and impacted the calculation process. Therefore, instead of using RGB colour space. The analysis has been conducted on other colour spaces such as HSV and YCbCr to identify the actual colour values of the pineapple skin, which is known as colour feature.

Obtain Image	Range Values of RGB Colour Components	Colour Distribution	Generated Result
 (a)			

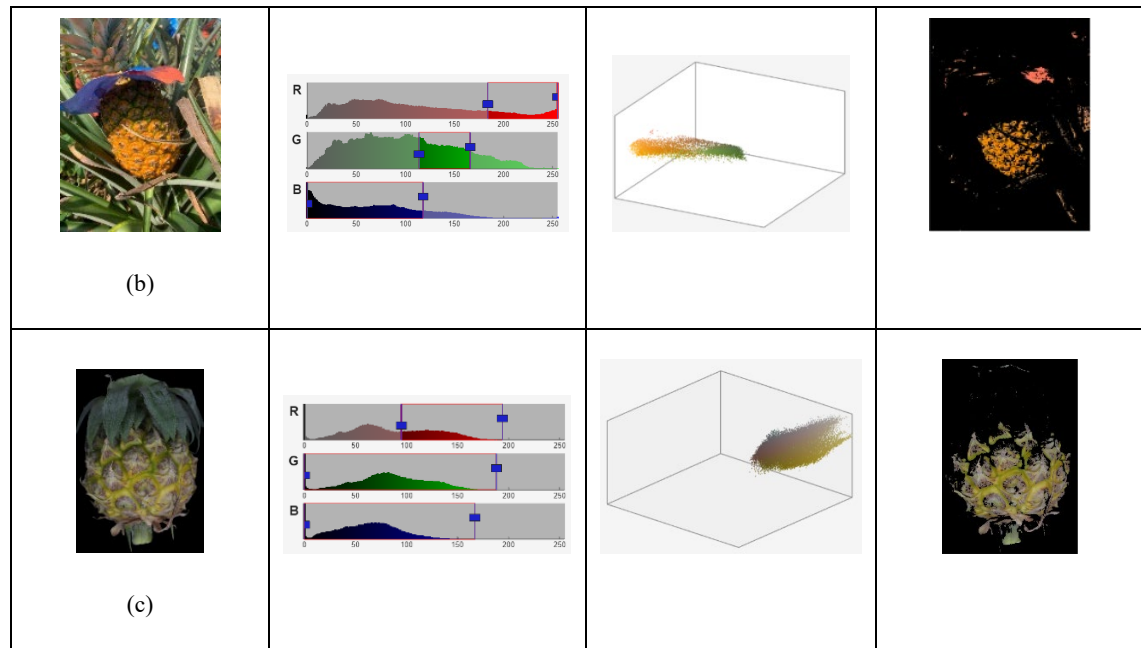


Fig. 4. Colour features analysis using RGB colour space

Fig. 4 presents the findings of the analysis, demonstrating that the RGB component values exhibit significant variation, which is essential for determining the colour value of each pixel in the pineapple skin. However, accurately identifying the true colour value remains challenging due to the wide range of RGB values, which influences the noise elimination process and may result in the unintended removal of key features within the ROI. For instance, in image (a), to effectively exclude the pineapple crown from the ROI, the green component value must be constrained between >50 and ≤ 250 , while the blue component must range from 0 to <150 . Nevertheless, the results indicate that the noise elimination process not only removes the pineapple crown but also eliminates regions within the ROI that share similar colour properties, potentially affecting the accuracy of feature extraction. A similar issue is observed in images (b) and (c), where portions of the ROI have also been inadvertently removed, as evidenced in the generated results section.

Next, the analysis is continued with the HSV colour space as shown in Fig. 5. The saturation (S) and value (V) components have a normalised range of 0 to 1, as shown in the figure. When compared to the RGB counterparts, the range of values for these components is quite narrow. Furthermore, the hue (H) component regularly has a threshold value below 1. As a result, the noise is successfully eliminated from the ROI as presents in generated result for image (a) and (c). However, the consequent image (b) is not completely clean, as some noise artifacts remain. Nevertheless, the ROI in image (b) exhibits more apparent clarity and focus compared to the ROI in image (b) from Fig. 4.

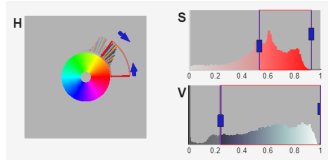
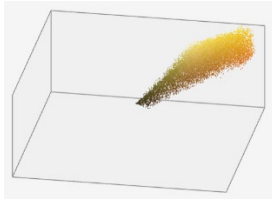
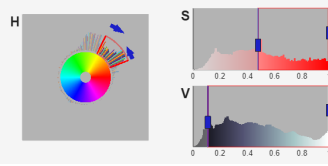
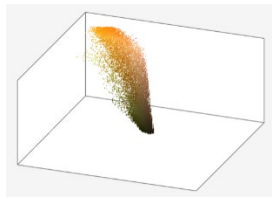
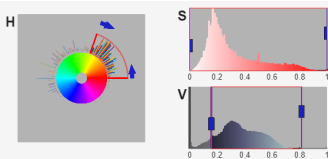
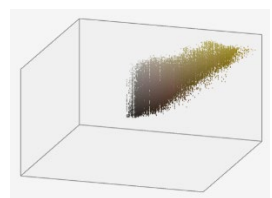
Obtain Image	Range Values of HSV Colour Components	Colour Distribution	Generated Result
 (a)			
 (b)			
 (c)			

Fig. 5. Colour features analysis using HSV colour space

The analysis is then carried out using the YCbCr colour space, as seen in Fig. 6. Comparable to the RGB colour system, the range value of each component in the YCbCr colour space is 0 to 255. Furthermore, the obtained result for each image is similar to that shown in Fig. 4. However, the generated result of image (b) in Fig. 6 has less noise than the generated image (b) in Fig. 4. Moreover, in general the range value of Y component is >10 and ≤ 230 . Then, the range value of Cb component is >0 and ≤ 150 . Next, the range value of Cr component is ≥ 140 and ≤ 180 .

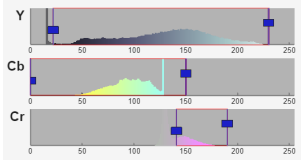
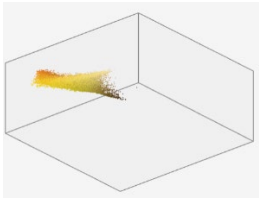
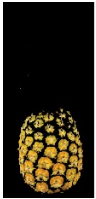
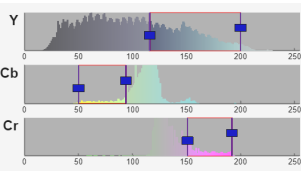
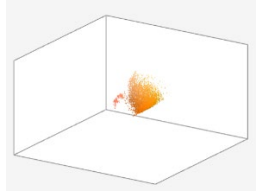
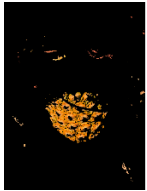
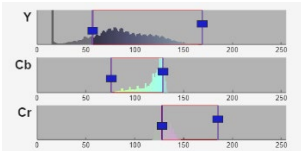
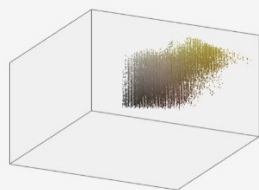
Obtain Image	Range Values of YCbCr Colour Components	Colour Distribution	Generated Result
 (a)			
 (b)			
 (c)			

Fig. 6. Colour features analysis using YCbCr colour space

Based on the comprehensive discussion, there is no specific value can be determined as the value of the pineapple skin colour. But the finding has shown that the range of the colour value from each component is possible to be used to extract the colour feature of the pineapple skin. Table 3 emphasises the finding from the conducted analysis.

Table 3. Range of the colour value for each component

R	G	B
≥ 100	> 50 and ≤ 250	0 to < 150
H	S	V
< 1	> 0 and < 1	> 0 and < 1
Y	Cb	Cr
> 10 and ≤ 230	> 0 and ≤ 150	≥ 140 and ≤ 180

Although the range of the colour value is shown in Table 3, still, further study needs to be conducted to ensure the extraction will extract the accurate ROI features. In other words, there is a possibility to discover the actual range or colour value of the pineapple skin.

4. FINDINGS

Following extensive analysis, the study revealed numerous important concerns that require action. First, there is concern about the pineapple image dataset's quality and representation. Second, the extraction approach used in the study has limitations that must be addressed. Finally, the classification approach must be refined to improve accuracy and dependability.

Data quality is the primary concern requiring improvement. High-quality images can be achieved by implementing a standard operating procedure (SOP) and selecting an appropriate camera. Key variables to consider when developing the SOP include lighting conditions, camera-to-object distance, minimal or controlled background noise, and saving images in PNG or JPEG file formats.

Subsequent analysis indicates that the HSV colour space is more reliable for extracting colour features from pineapple skin. However, post-extraction image enhancement is necessary to improve output quality and minimise classification errors. The recommended threshold values for each HSV component, as presented in Table 3, can be applied for this purpose.

Finally, the classification approach requires refinement to enhance the accuracy. The algorithm should be enhanced to incorporate two critical parameters in its computations: colour features and pineapple surface morphology. Furthermore, the logical operators employed in the classification process must be optimised to ensure accurate categorical assignment of images.

The identified research gaps present significant opportunities for further investigation and methodological advancement in this domain. Empirical validation is necessary, with systematic verification of experimental results to confirm the reliability of the proposed solution.

5. CONCLUSION AND FUTURE WORK

In conclusion, environmental interference in pineapple imaging represents a critical challenge requiring resolution. Image noise has been identified as a primary contributor to classification inaccuracies. Analytical findings reveal two key research priorities: first, the need to precisely define pineapple skin colour characteristics, and second, the development of robust classification methodologies. This study establishes HSV colour space as the most effective model for extracting region-of-interest colour features from pineapple skin. Furthermore, this study has proposed incorporating morphological parameters (specifically surface area measurements) and optimised logical operators within the classification algorithm to address these challenges.

The identified research gaps, namely colour definition and classification refinement, are fundamental challenges that must be addressed to enable the development of reliable autonomous pineapple recognition systems. Future research directions should prioritise experimental validation to demonstrate the practical efficacy of the proposed methodology. Additionally, this investigation will explore the applicability of the YOLOv7 algorithm for processing suboptimal pineapple images under real-world conditions. Specifically, subsequent work will focus on developing the algorithm's capability to assess pineapple maturity directly on the tree, thereby facilitating automated harvest determination.

6. ACKNOWLEDGMENTS/FUNDING

We are grateful to UiTM Johor Branch Pasir Gudang Campus and UiTM Terengganu Branch Kuala Terengganu Campus for providing us with the resources and facilities we require to further our academic goals. We thank our department for creating a very efficient research environment and acknowledge all individuals who contributed to this endeavour.

7. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

8. AUTHORS' CONTRIBUTIONS

Mohamad Faizal Ab Jabal: conceptualisation, writing, validation, and data curation. **Muhammad Irfan Rozlan:** concept, approach, and software. **Azrina Suhaimi:** handles conceptualization, methodology, formal analysis, investigation, and proofreading. **Harshida Hasmy:** manages resources, formal analysis, and investigations.

9. REFERENCES

- Arboleda, E., Jesus, C., & Tia, L. (2021). Pineapple maturity classifier using image processing and fuzzy logic. *IAES International Journal of Artificial Intelligence (IJ-AI)*, 10(4), 830–838. <https://doi.org/10.11591/ijai.v10.i4.pp830-838>
- Cai, J., Tao, J., Ma, Y., Fan, X., & Cheng, L. (2020). Fruit image recognition and classification method based on improved single shot multi-box detectors. *Journal of Physics: Conference Series*, 1629(1), 012010. <https://doi.org/10.1088/1742-6596/1629/1/012010>
- Chang, H., Meng, Q., Wu, Z., Tang, L., Qiu, Z., Ni, C., Chu, J., Fang, J., Huang, Y., & Li, Y. (2025). Accurate ripening stage classification of pineapple based on a visible and near-infrared hyperspectral imaging system. *Journal of AOAC International*, 108(3), 293–303. <https://doi.org/10.1093/jaoacint/qsaf010>
- Chen, Y., Zheng, L., & Peng, H. (2023). Assessing pineapple maturity in complex scenarios using an improved RetinaNet algorithm. *Engenharia Agricola*, 43(2), e20220180. <https://doi.org/10.1590/1809-4430-eng.agric.v43n2e20220180/2023>
- KSF Enterprise MD2. (2024). *Pembekal benih nanas MD2* [Facebook page]. Facebook. https://www.facebook.com/Ksfenterprisemd2/photos?locale=ms_MY
- Jabal, M. F. A., Hamid, S., Othman, N. Z. S., & Rahim, M. S. M. (2022). Spot filtering adaptive thresholding (SFAT) method for early pigment spot detection on iris surface. In Saba, T., Rehman, A., Roy, S. (Eds), *Prognostic models in healthcare: AI and statistical approaches* (pp. 347–370). Springer Nature Singapore. https://doi.org/10.1007/978-981-19-2057-8_13
- Jabal, M. F. A., Abdullah, A. N. H., Najjar, F. H., Hamid, S., Khalid, A. K., & Manan, W. D. W. A. (2024). A novel color feature for the improvement of pigment spot extraction in iris images. *Journal of Image and Graphics*, 12(4), 410–416. <https://doi.org/10.18178/joig.12.4.410-416>
- Lai, Y., Ma, R., Chen, Y., Wan, T., Jiao, R., & He, H. (2023). A pineapple target detection method in a field environment based on improved YOLOv7. *Applied Sciences*, 13(4), 2691. <https://doi.org/10.24191/jcrim.v10i2.546>

<https://doi.org/10.3390/app13042691>

- Li, J., Liu, Y., Li, C., Luo, Q., & Lu, J. (2024). Pineapple detection with YOLOv7-tiny network model improved via pruning and a lightweight backbone sub-network. *Remote Sensing*, 16(15), 2805. <https://doi.org/10.3390/rs16152805>
- Malaysian Pineapple Industry Board. (2024). *Home* [Official government website]. <https://www.mpib.gov.my/en/home/>
- Manik, F. Y., Harumy, T. H. F., Akasah, W., Hidayat, W., Simanjuntak, R. F., & Sianturi, V. J. (2024, April 19). Identification of pineapple maturity utilizing digital image using hybrid machine learning method. *AIP Conference Proceedings*, 2987(1), 020050. <https://doi.org/10.1063/5.0199826>
- Mukhiddinov, M., Muminov, A., & Cho, J. (2022). Improved classification approach for fruits and vegetables freshness based on deep learning. *Sensors*, 22(21), 8192. <https://doi.org/10.3390/s22218192>
- Roslan, M. I., Ibrahim, Z., Adnan, N. A. K., Diah, N. M., Narawi, N. A. A., & Arif, Y. M. (2023). Fruit detection and recognition using Faster R-CNN with FPN30 pre-trained network. In *2023 IEEE 8th International Conference on Recent Advances and Innovations in Engineering (ICRAIE)* (pp. 1–6). IEEE. <https://doi.org/10.1109/ICRAIE59459.2023.10468169>
- Shiu, Y.-S., Lee, R.-Y., & Chang, Y.-C. (2023). Pineapples' detection and segmentation based on Faster and Mask R-CNN in UAV imagery. *Remote Sensing*, 15(3), 814. <https://doi.org/10.3390/rs15030814>
- Sujitra, A. (2024). *Phulae pineapple translucent flesh symptom dataset* [Dataset]. Kaggle. <https://www.kaggle.com/datasets/sujitraarw/phulae-pineapple-translucency-defect>
- Trinh, T. H., & Nguyen, H. H. C. (2023). Implementation of YOLOv5 for real-time maturity detection and identification of pineapples. *Traitement du Signal*, 40(4), 1445–1455. <https://doi.org/10.18280/ts.400413>
- Syazwani, R. W. N., Asraf, H. M., Amin, M. M. S., & Dalila, K. N. (2022). Automated image identification, detection and fruit counting of top-view pineapple crown using machine learning. *Alexandria Engineering Journal*, 61(2), 1265–1276. <https://doi.org/10.1016/j.aej.2021.06.053>
- Wang, H. H., & Chai, S. Y. (2022). Novel feature extraction for pineapple ripeness classification. *Journal of Telecommunications and Information Technology*, 1, 14–22. <https://doi.org/10.26636/jtit.2022.156021>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Thyroid Insight: Navigating Disease Data Through Interactive Visualization with Prediction

Mohd Nizam Osman^{1*}, Azim Md Nasib², Khairul Anwar Sedek³, Nor Arzami Othman⁴, Mushahadah Maghribi⁵

^{1,2,3,4}Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Perlis Branch, Arau Campus, 02600 Arau, Perlis, Malaysia.

⁵Politeknik Tuanku Syed Sirajuddin, 02600 Arau, Perlis, Malaysia.

ARTICLE INFO	ABSTRACT
<i>Article history:</i> Received 30 June 2025 Revised 15 August 2025 Accepted 15 August 2025 Published 1 September 2025 <i>Keywords:</i> Thyroid Disorders Data Visualization Predictive Modelling Random Forest Interactive Dashboard TAM <i>DOI:</i> 10.24191/jcrinn.v10i2	The thyroid gland, located in the neck, plays a crucial role in regulating metabolism, growth, and energy through hormones such as thyroxine (T4) and triiodothyronine (T3). Disorders such as hypothyroidism, hyperthyroidism, and thyroid cancers are often linked to iodine deficiency and genetics. However, limited public awareness and delayed diagnosis can lead to severe health complications. Analysing thyroid disease data is challenging due to its complexity and unstructured nature, making advanced analytical techniques essential. This paper addresses these challenges by developing an interactive dashboard with predictive capabilities. The system integrates Big Data analytics and predictive modelling to improve understanding and support proactive management of thyroid health. It follows a structured methodology, including planning, analysis, design, development, and testing, using data from Kaggle and the UCI Machine Learning Repository. The dashboard employs Microsoft Power BI for visualizations and the Random Forest algorithm for predictive modelling. Evaluation using the Technology Acceptance Model (TAM) with 35 respondents produced encouraging results across dimensions such as Perceived Ease of Use (4.28), Perceived Usefulness (4.61), Attitude (4.54), and Intention to Use (4.50). User feedback highlighted the dashboard's intuitive design, clarity in presenting complex information, and potential to raise awareness about thyroid health. While the findings are based on a limited evaluation, results indicate that the system may contribute to improving public awareness, supporting early detection, and empowering users to make informed health decisions. With future improvements, such as real-time data integration and expanded datasets, the system could further enhance healthcare practices and public education regarding thyroid diseases, promoting proactive health management.

1. INTRODUCTION

Technology has rapidly transformed many aspects of daily life, particularly in the healthcare sector, where advances in data analytics and visualization are revolutionizing how medical information is interpreted and applied (Wang et al., 2024). One such advancement is the use of Big Data in managing health conditions such as thyroid disorders. These disorders, including hypothyroidism, hyperthyroidism, and thyroid cancer, affect millions of individuals worldwide, creating a growing need for effective diagnosis and management strategies (Keestra et al., 2021). However, despite their prevalence, early detection and awareness remain limited due to the complexity of thyroid health data and a lack of accessible tools for the general public (Zhao et al., 2024).

Many thyroid patients face challenges in understanding their condition and accessing appropriate resources. According to Fariduddin et al. (2024), integrating predictive technologies with healthcare

^{1*} Corresponding author. E-mail address: mohdnizam@uitm.edu.my

systems can improve early detection and patient awareness. While previous studies have investigated medical data management systems, few have successfully combined Big Data analytics with real-time health monitoring for thyroid disorders. Some initiatives have developed interactive dashboards that provide users with access to their health data, symptoms tracking, and potential treatment options. However, these systems often lack advanced predictive models and real-time monitoring capabilities, which are crucial for proactive thyroid health management.

This study builds on prior work by integrating Big Data analytics, predictive modelling, and interactive visualization to develop the "Thyroid Insight" dashboard. The system aims to raise public awareness of thyroid diseases, support early detection, and enable users to take a more proactive role in managing their health. Utilizing data from reputable sources such as Kaggle and the UCI Machine Learning Repository, and applying predictive algorithms such as Random Forest, the dashboard forecasts potential health risks and offers actionable insights (Muniswamaiah et al., 2023). By combining predictive analytics with real-time data collection tools, the system offers a more comprehensive solution that bridges the gap between healthcare providers and individuals.

The purpose of this study was to evaluate the effectiveness of the dashboard in supporting early detection and delivering actionable insights into thyroid health. By assessing user satisfaction, system usability, and overall functionality, the project aims to provide a practical tool for improving thyroid disease management and fostering a more informed, health-conscious society. The "Thyroid Insight" system was developed using Microsoft Power BI for visualization and the Random Forest algorithm for predictive analytics, resulting in an intuitive platform designed to enhance the understanding and management of thyroid health.

2. RELATED WORK

In the field of thyroid disease management, the application of big data analytics and machine learning has shown great promise for improving diagnosis, treatment, and public awareness. Thyroid disorders such as hypothyroidism, hyperthyroidism, and thyroid cancers affect millions worldwide. Besides, the thyroid gland is essential for regulating metabolism through the production of hormones such as T4 and T3, and disruptions in this system can lead to various health complications. Managing these conditions involves interpreting complex datasets, including hormone levels, patient records, and medical histories, which makes traditional diagnostic methods less effective (Zhao et al., 2024). Advanced analytical tools and visualization platforms are therefore essential for addressing this complexity.

Numerous studies have explored predictive modelling and data visualization for thyroid health, but most focus on either predictive accuracy or data presentation, rather than integrating both into a user-accessible, real-time platform. Several relevant works have made strides in these areas, providing a foundation for advancing visualization and predictive systems in thyroid disease management.

2.1 CancerMAS dashboard

The CancerMAS Dashboard, developed by the Advanced Analytics Engineering Centre at Universiti Teknologi MARA, is a web-based platform for visualizing cancer incidence data in Malaysia (Hamidi et al., 2020). It presents data on cancer types, gender, ethnicity, and geographical distribution using bar charts, scatter plots, and heat maps. While it demonstrates strong visualization capabilities for public health analysis, it lacks predictive modelling, personalized risk assessment, and application to thyroid disease. Our work adapts these visualization strengths while integrating machine learning-based prediction tailored to thyroid disorders.

2.2 Thyroid disease prediction using selective features and machine learning techniques

Chaganti et al. (2022) applied selective feature selection, to improve thyroid disease prediction accuracy. Using Random Forest, they achieved 0.99 accuracy, demonstrating the value of targeted feature selection. However, their system remains model-centric, without translating results into visuals, interactive

tools for non-specialist users. In contrast, our approach embeds predictive models within a user-friendly dashboard, enabling both technical accuracy and public accessibility.

2.3 Detecting thyroid disease using optimized machine learning model based on differential evolution

Gupta et al. (2024) developed an optimized AdaBoost model enhanced with Differential Evolution (DE) and CTGAN for data augmentation, achieving 0.998 accuracy. While this approach addresses data imbalance and achieves high performance, it does not include a visualization interface for communicating predictions to patients or healthcare providers. Our system builds on such predictive strength by combining it with interactive visualization and a focus on user engagement.

2.4 International agency for research on cancer

The International Agency for Research on Cancer (IARC) platform provides global cancer statistics, including thyroid cancer data, via interactive heatmaps (Terrasse, 2025) . Users can filter by demographic variables and explore trends over time. While highly customizable for epidemiological analysis, it does not offer predictive capabilities or personalized health insights. Our dashboard addresses this gap by merging population-level visualization with individual-level predictions for thyroid health.

To position this research within the existing literature, Table 1 presents a comparative overview of the related works discussed. The table evaluates each system based on key attributes, including visualization capabilities, predictive modelling, real-time data integration, and focus area, while also noting their primary limitations. This structured comparison reveals a recurring gap: existing works often excel in either visualization or prediction but rarely combine both within a single, user-accessible platform tailored for thyroid disease management. Moreover, none of the reviewed studies fully integrate real-time monitoring, despite its growing importance for proactive health management. The proposed “Thyroid Insight” dashboard addresses these shortcomings by unifying predictive analytics, interactive visualization, and the potential for real-time data integration, thereby offering a more comprehensive solution to enhance public awareness and support early detection.

Table 1. Comparison table of related work

Study/System	Visualization	Prediction	Real-Time Data	Focus Area	Gap Filled by This Study
CancerMAS Dashboard (Hamidi et al. (2020)	✓ Strong charts, maps	✗ None	✗ No	No prediction or thyroid-specific data	Combines similar visualization with thyroid-specific predictions
Chaganti et al. (2022))	✗ None	✓ RF (0.99 accuracy)	✗ No	No interactive/public-friendly visualization	Embeds prediction into user-friendly dashboard
Gupta et al. (2024)	✗ None	✓ AdaBoost+DE (0.998)	✗ No	No visualization or public engagement	Combines high-accuracy prediction with visualization
IARC (Terrasse, 2025)	✓ Interactive heatmaps	✗ None	✗ No	No predictive or personalized insights	Adds predictive analytics to population-level visualizations

3. METHODOLOGY

This research adopted a structured System Development Life Cycle (SDLC) approach to guide the development of the proposed system. The SDLC methodology provides a systematic framework that ensures progress through clearly defined stages, emphasizing comprehensive analysis, thoughtful design, efficient development, and rigorous testing. The phases encompassed in this methodology are Planning, Analysis, Design, Development, and Testing, which collectively contributed to the creation of an interactive dashboard for thyroid disease management with predictive capabilities. A breakdown of each phase is presented below.

3.1 Planning phase

The Planning phase is the foundation of the project, where the problem statement is defined and the research objectives, scope, and significance are established. This phase is critical for setting the direction for the entire study. To achieve this, a comprehensive literature review was conducted to identify current challenges in thyroid disease prediction and management. The primary objective was to integrate Big Data analytics with predictive modelling to enhance public awareness, support early detection, and improve understanding of thyroid diseases. The project's scope was defined as the development of an interactive dashboard using appropriate datasets. The datasets used in this study were obtained from the UCI Machine Learning Repository (Quinlan, 1987) and Kaggle's Thyroid Disease dataset (Kaggle, 2021). These sources were selected for their public availability, structured format, and widespread use in peer-reviewed studies, making them reliable benchmarks for machine learning research (Fernández-Delgado et al., 2014). The UCI dataset contains patient records with hormone measurements, demographic details, and diagnostic labels derived from clinical examinations, while the Kaggle dataset includes more recent and diverse cases contributed by the medical data science community.

3.2 Analysis phase

The Analysis phase focuses on identifying and examining the system requirements, relevant literature, and technologies that align with the project's objectives. Key components analysed in this phase include Big Data sources, such as patient records and hormone level measurements, data visualization tools, such as Microsoft Power BI, and Machine Learning Models, particularly the Random Forest algorithm. A comprehensive review of existing research was also conducted to evaluate current solutions in thyroid disease prediction and related healthcare applications. In addition, this phase involved assessing the hardware and software requirements necessary for effective system implementation. The selected technologies include Microsoft Power BI for visualization, Python and Jupyter Notebook for machine learning model development, and the Random Forest algorithm for predictive analysis. The primary goal of the analysis phase was to ensure that the proposed system would meet performance and user expectations while maintaining compatibility with the chosen tools and resources.

3.3 Design phase

The Design phase defines the system architecture, user interface, and data structure. This phase involved creating visual representations, such as Entity Relationship Diagrams (ERD), to conceptualize data flow and system interactions. A sitemap was developed to illustrate user navigation within the interactive dashboard. The interface incorporated key features, including data visualization elements such as graphs, pie charts, and prediction input forms. The dashboard was designed to be intuitive and user-friendly, enabling users to easily interpret complex thyroid health data. Wireframes were created for the main section of the dashboard, including the homepage, disease statistics, thyroid disorder types, and the prediction result page. These design activities were essential to ensuring a coherent layout, efficient navigation, and an optimal user experience.

3.4 Development phase

The Development phase involved building the system based on the designs established in the previous phase. Two primary activities were carried out during this phase:

- i. **Developing the Web Application:** The system was implemented as a web-based application using Python and Microsoft Power BI as the primary development tools. Python was employed for

backend scripting, particularly to handle machine learning models, data processing, and prediction tasks. Microsoft Power BI was used to create interactive data visualizations for displaying thyroid disease statistics. The application was designed to retrieve data from various sources, perform data cleaning and preprocessing, and present the processed information in an accessible format through the dashboard interface.

- ii. **Integration of Predictive Models:** After the core web application was built, the next step was to integrate the predictive models. The Random Forest algorithm was applied to develop predictive models for thyroid disease classification and prediction based on patient data, including hormone levels and medical history. Data preprocessing, feature selection, and handling of missing values were essential steps before training the models. The Random Forest algorithm (Breiman, 2001) was employed for thyroid disease classification and risk prediction. Hyperparameter tuning was conducted using Grid Search with 5-fold cross-validation to determine the optimal number of trees (`n_estimators`), maximum tree depth (`max_depth`), and the number of features considered at each split (`max_features`). To mitigate overfitting, the minimum samples required to split a node (`min_samples_split`) was adjusted, and tree depth was constrained. Model evaluation was performed using accuracy, precision, recall, and F1-score metrics (Sokolova & Lapalme, 2009), with the dataset split into 80% training and 20% testing subsets. The final tuned model achieved an accuracy of 98.5% and an F1-score of 0.984, outperforming baseline models such as Logistic Regression and Decision Tree. Once validated, the model was serialized using the `joblib` library and integrated into the application, enabling real-time predictions from user-provided inputs, including hormone levels (T3, T4, TSH), patient age, and medical history.

3.5 Testing phase

The Testing phase was carried out to evaluate the system's functionality, usability, and reliability. Usability testing was conducted with a group of 35 participants from diverse backgrounds to ensure accessibility and inclusivity. A Technology Acceptance Model (TAM) survey was administered to measure user acceptance focusing on key dimensions such as Perceived Ease of Use, Perceived Usefulness, Attitude, and Intention to Use. Participant Feedback played a crucial role in refining the system, enhancing the user interface, and ensuring that the final product effectively addressed the needs of the target audience.

4. RESULTS AND DISCUSSIONS

4.1 Interfaces of the Thyroid Insight system

This section presents the interfaces and functionality of the proposed system. Designed with user accessibility as a priority, the system features an intuitive interface suitable for both healthcare professionals and the general public. The main page includes several key sections that provide access to core features, such as interactive data visualizations, disease prediction results, and detailed information on thyroid health. Fig. 1 illustrates the system's main interface.

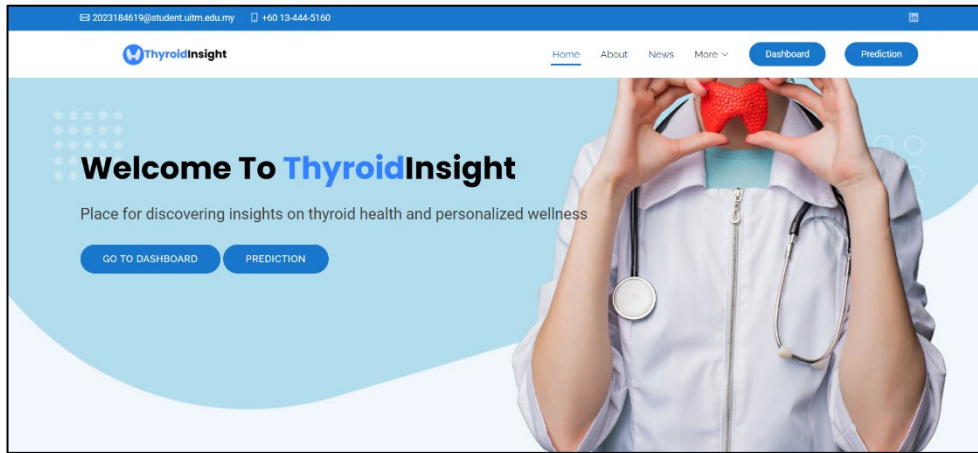


Fig. 1. Main interface of Thyroid Insight system

The main interface of the system is the dashboard, which provides an overview of thyroid disease data. It features interactive visualizations, including graphs, pie charts, and trend lines developed using Microsoft Power BI. These visualization enables users to explore thyroid disease data through various filters and viewing options. The dashboard presents key insights such as diseases prevalence, age distributions, and variations in hormone levels, thereby supporting informed health decision-making. Fig.2 – Fig. 7 present the interfaces of visualization dashboards.

The interface also features a prediction results section, where users can input their health data, such as age, hormone levels, and medical history, to receive personalized risk assessments for thyroid diseases. These predictions are generated using machine learning models, specifically the Random Forest algorithm, trained on dataset from the UCI Machine Learning Repository. Integrating these models into the system enables the delivery of real-time, personalized health insights, thereby supporting proactive thyroid disease management.

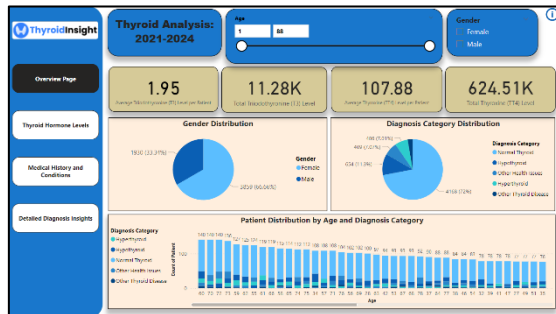


Fig. 2. Overview page

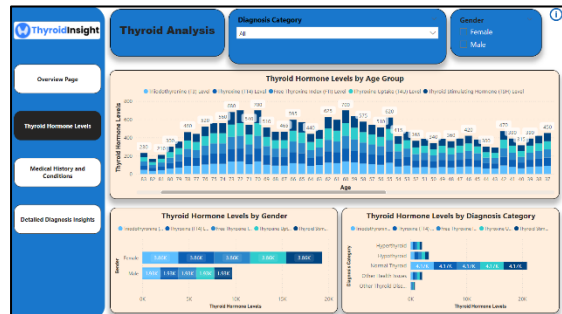


Fig. 3. Thyroid hormone levels page

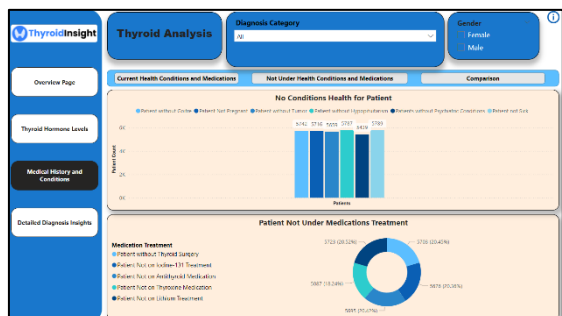
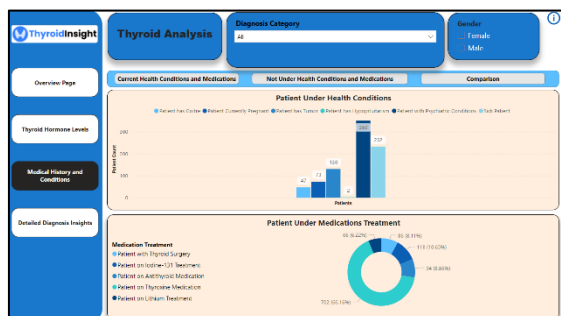


Fig. 4. Medical history page

Fig. 6. Comparison page

Fig. 5. Medical history page

Fig. 7. Detailed diagnosis page

The prediction module is a key feature of the Thyroid Insight system. Upon entering their health data, users receive predictive results regarding their thyroid health status. This feature leverages the Random Forest algorithm, which analyses multiple data points and generates predictions based on established patterns within the dataset. Fig. 8 illustrates the interfaces for the prediction form.

The prediction results are displayed on a dedicated page, as shown in Fig. 9, presenting users with their estimated likelihood of developing a thyroid disorder. For instance, the system may indicate a user’s risk of hypothyroidism based on their hormone levels, age, and medical history. These predictions are designed to empower individuals to take proactive measures in maintaining their thyroid health. In addition, the results include tailored recommendations for managing health, informed by the specific risk factors identified by the model.

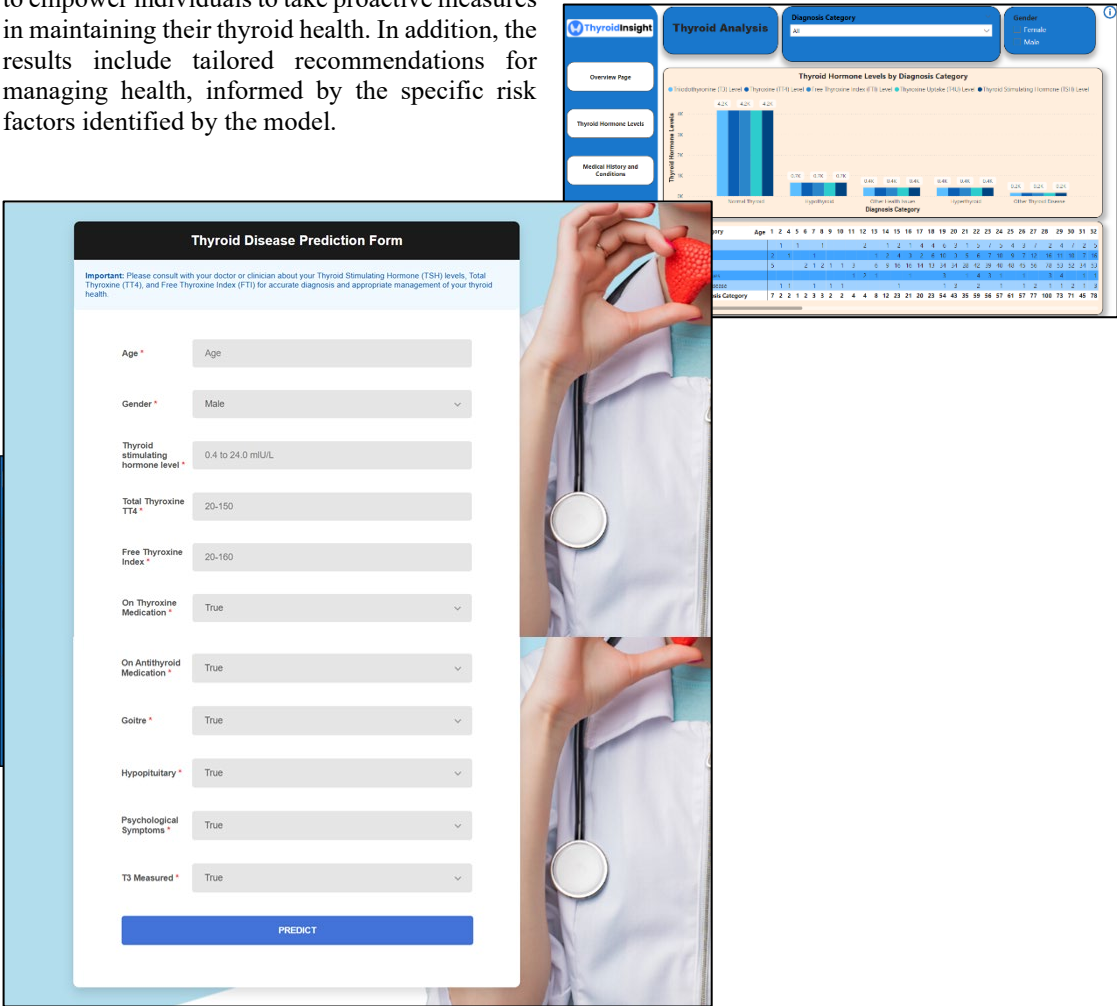


Fig. 8. Prediction form

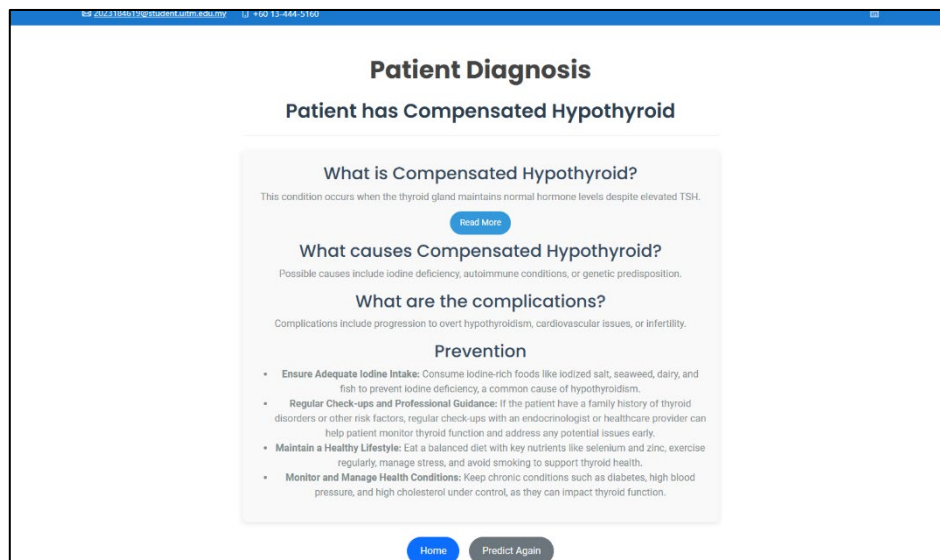


Fig. 9.

Prediction result page

4.2 The Findings of The Study

This section presents the dimensions assessed during the user acceptance testing of the Thyroid Insight system. The evaluation employed the Technology Acceptance Model (TAM), which measures four key dimensions which are Perceived Ease of Use (PEU), Perceived Usefulness (PU), Attitude (ATT), and Intention to Use (BI). Each dimension reflects a distinct aspect of user experience and offers valuable insights into how the system was received by participants. The results from the user testing were compared across these dimensions to highlight areas of strength and identify opportunities for improvement. Fig. 10 presents a comparison of the average scores for PEU, PU, ATT, and BI.

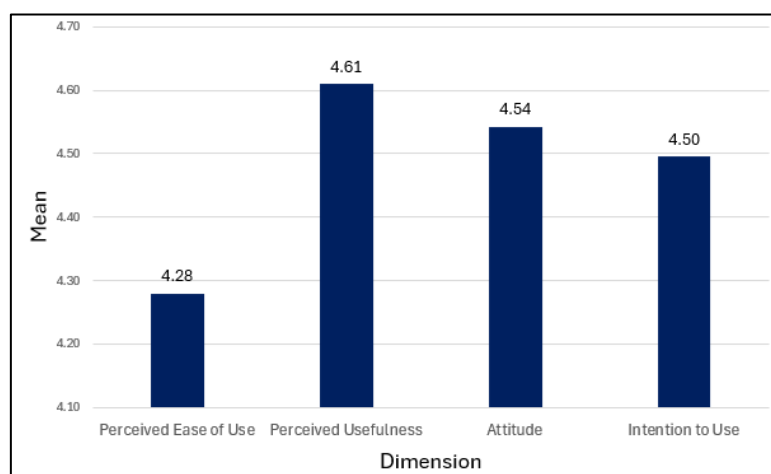


Fig. 10. Comparison of each dimension

As shown in Fig. 10, the highest-rated dimension is Perceived Usefulness (PU), with an average score of 4.61. This suggests that respondents view the "Thyroid Insight" dashboard as highly beneficial and effective in enhancing their understanding of thyroid diseases. The Attitude (ATT) dimension follows closely, with an average score of 4.54, reflecting a strong positive attitude among users and indicating their satisfaction, enjoyment, and appreciation of the dashboard's features and functionality. Intention to Use

(BI) ranks next, with an average score of 4.50, demonstrating that most respondents intend to continue using the dashboard and recommend it to others, further underscoring its perceived value and impact.

Perceived Ease of Use (PEU) has the lowest average score among the dimensions, at 4.28. While still positive, this slightly lower score suggests that some users encountered minor challenges with the dashboard's usability. Feedback from the usability testing revealed that certain participants, particularly those less familiar with technology, found the data input process and navigation between dashboard sections less intuitive. Contributing factors may include the number of required input fields, the presence of technical terminology in the prediction forms, and the absence of step-by-step prompts or contextual help. The implication for future development is that simplifying the navigation structure, reducing cognitive load, and integrating supportive features such as tooltips, guided workflows, or onboarding tutorials could further improve PEU. Enhancing ease of use is not only expected to increase satisfaction but may also positively influence both attitude and intention to use, as perceived simplicity plays a critical role in TAM-based adoption models.

Overall, the Thyroid Insight system received positive feedback across all four TAM dimensions, particularly in terms of Perceived Usefulness and Intention to Use. The system was recognized for its ability to enhance understanding and awareness of thyroid diseases, offering users a valuable tool for early detection and health management. While the system scored well on Perceived Ease of Use, there is room for improvement in simplifying the user input process and providing more guidance for less tech-savvy users. Attitude results suggest that while users are generally satisfied, there is potential for further customization and features to increase engagement and usefulness. By addressing these areas of improvement, future iterations of the system can enhance both user experience and its overall impact on thyroid disease management.

5. CONCLUSION

The development of the Thyroid Insight system represents a significant contribution to the field of thyroid disease management through the integration of Big Data analytics, machine learning models, and advanced data visualization tools. This research addressed the need for a user-friendly, interactive platform that promotes public awareness, supports early detection, and enhances understanding of thyroid health. By analysing large datasets from reliable sources, the system delivers insightful predictions and visualizations that enable users to manage their thyroid health more effectively.

The integration of predictive models, such as the Random Forest algorithm, enabled the system to provide personalized risk assessments based on individual health data, making it a valuable tool for early intervention and proactive health management. The system's intuitive dashboard, developed using Microsoft Power BI, delivers an engaging and informative user experience, presenting complex medical data in clear and accessible format for both healthcare professionals and the general public.

Furthermore, this research contributed to the advancement of usability testing methodologies by engaging a diverse group of 35 participants in a comprehensive evaluation of the system. The positive feedback and high satisfaction levels validated the system's effectiveness, confirming that it met its objectives of improving understanding and raising awareness about thyroid diseases. The usability testing also identified areas for enhancement, including simplifying the prediction input process and introducing additional customization features, which will inform future system improvements.

The integration of Big Data and machine learning to predict thyroid disease risk represents an innovative approach to healthcare management. The system's ability to deliver personalized insights empowers individuals to take a more active role in managing their health, while also supporting healthcare providers in delivering targeted interventions. By addressing the complexity of thyroid disease data, the system fosters a more informed and proactive approach to thyroid health management.

Overall, the Thyroid Insight system provides a practical and effective solution for understanding and managing thyroid diseases. Its success demonstrates the potential of integrating advanced technologies to improve healthcare outcomes, offering a valuable tool for individuals seeking to monitor and manage their thyroid health. The contributions of this research pave the way for future advancements in healthcare

technology, particularly through the incorporation of more personalized dataset and real-time health monitoring, which could further enhance the system's accuracy, usability, and overall impact.

6. ACKNOWLEDGEMENTS/FUNDING

The authors would like to acknowledge the support of Universiti Teknologi MARA Perlis Branch, and the Faculty of Computer and Mathematical Sciences for providing the facilities and resources necessary to conduct this research.

7. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interest.

8. AUTHORS' CONTRIBUTIONS

Mohd Nizam Osman and **Azim Md Nasib** conceived of the original and presented idea and developed the theory as well as the system development process. **Khairul Anwar Sedek**, **Mushahadah Maghribi** and **Nor Arzami Othman** verified the analytical methods. Mohd Nizam Osman encouraged Azim Md Nasib to explore a specific aspect, supervised the findings, and conducted the experiments. Mushahadah Maghribi contributed to the interpretation of the results. All authors discussed the results and contributed to the final version of the manuscript.

9. REFERENCES

- Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5–32. <https://doi.org/10.1023/A:1010933404324>
- Chaganti, R., Rustam, F., De La Torre Díez, I., Mazón, J. L. V., Rodríguez, C. L., & Ashraf, I. (2022). Thyroid disease prediction using selective features and machine learning techniques. *Cancers*, 14(16), 3914. <https://doi.org/10.3390/cancers14163914>
- Fariduddin, M. M., Haq, N., & Bansal, N. (2024). *Hypothyroid myopathy*. StatPearls Publishing. <https://www.ncbi.nlm.nih.gov/books/NBK519513/>
- Fernández-Delgado, M., Cernadas, E., Barro, S., & Amorim, D. (2014). Do we need hundreds of classifiers to solve real world classification problems? *Journal of Machine Learning Research*, 15, 3133–3181.
- Gupta, P., Rustam, F., Kanwal, K., Aljedaani, W., Alfarhood, S., Safran, M., & Ashraf, I. (2024). Detecting thyroid disease using optimized machine learning model based on differential evolution. *The International Journal of Computational Intelligence Systems*, 17(1). <https://doi.org/10.1007/s44196-023-00388-2>
- Hamidi, S. R., Lokman, A. M., Shuhidan, S. M., & Hilmi, M. I. M. N. (2020). CancerMAS dashboard: Data visualization of cancer cases in Malaysia. *Journal of Physics. Conference Series* 1529(2), 022019. <https://doi.org/10.1088/1742-6596/1529/2/022019>
- Kaggle. (2021). *Thyroid disease* [Dataset]. Kaggle. <https://www.kaggle.com/>
- Keestra, S., Tabor, V., & Alvergne, A. (2021). Reinterpreting patterns of variation in human thyroid

- function. *Evolution, Medicine and Public Health*, 9(1), 93–112. <https://doi.org/10.1093/emph/eoaa043>
- Kotsiantis, S. B., Kanellopoulos, D., & Pintelas, P. E. (2006). Data preprocessing for supervised learning. *International Journal of Computer Science*, 1(2), 111–1
- Muniswamaiah, M., Agerwala, T., & Tappert, C. C. (2023). Big data and data visualization challenges. In *2023 IEEE International Conference on Big Data 173 (BigData)* (pp. 6227-6229). <https://doi.org/10.1109/BigData59044.2023.10386491>
- Quinlan, J. R. (1987). Simplifying decision trees. *International Journal of Man-Machine Studies*, 27(3), 221–234. [https://doi.org/10.1016/S0020-7373\(87\)80053-6](https://doi.org/10.1016/S0020-7373(87)80053-6)
- Sokolova, M., & Lapalme, G. (2009). A systematic analysis of performance measures for classification tasks. *Information Processing & Management*, 45(4), 427–437. <https://doi.org/10.1016/j.ipm.2009.03.002>
- Terrasse, V. (2025). *Homepage – IARC*. World Health Organization. <https://www.iarc.who.int>
- Wang, X., Wu, Z., Liu, Y., Wu, C., Jiang, J., Hashimoto, K., & Zhou, X. (2024). The role of thyroid-stimulating hormone in regulating lipid metabolism: Implications for body–brain communication. *Neurobiology of Disease*, 106658. <https://doi.org/10.1016/j.nbd.2024.106658>
- Zhao, X., Fenggui, B., Caixia, L., & Jun-E, Z. (2024). Distress, illness perception and coping style among thyroid cancer patients after thyroidectomy: A crosssectional study. *European Journal of Oncology Nursing*, 69, 102517. <https://doi.org/10.1016/j.ejon.2024.102517>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Heart Failure Detection Using Scaled Conjugate Gradient Method and Naïve Bayes

Norpah Mahat^{1*}, Norazwana Saidin @ Zubir², Izleen Ibrahim³, Jasmani Bidin⁴,
Mohamad Najib Mohamad Fadzil⁵, Sharifah Fahriyah Syed Abas⁶, Siti Sarah
Raseli⁷

^{1,2,3,4,5,6,7}Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Perlis Branch, Arau Campus, 02600 Arau, Perlis, Malaysia.

ARTICLE INFO

Article history:

Received 20 June 2025

Revised 18 August 2025

Accepted 19 August 2025

Published 1 September 2025

Keywords:

Heart Failure

Artificial Neural Network

Scaled Conjugate Gradient

Naïve Bayes

DOI:

10.24191/jcrinn.v10i2.544

ABSTRACT

Heart failure known as high mortality rates is a serious pathophysiological condition characterized and substantial long-term healthcare costs. Early detection is crucial, as the disease tends to progress without timely and appropriate intervention. This study aims to predict the risk of heart failure using structured clinical data and to leverage deep learning techniques to enhance the accuracy of risk assessment. The core objective is to demonstrate that early identification of heart failure indicators can significantly improve patient outcomes, potentially distinguishing between life and death. Recognizing these early warning signs provides a better opportunity for preventive care and timely treatment. To achieve this, two algorithms were employed: the Scaled Conjugate Gradient method within an Artificial Neural Network (ANN) framework, and the Naïve Bayes classifier. A Feed-Forward Neural Network (FFNN) was utilized as the primary classifier to detect the presence of heart failure. The neural network architecture used in this study consisted of 12 input neurons, 20 hidden layers, and a single output layer. The performance results revealed that the ANN achieved an accuracy of 86.7%, while the Naïve Bayes classifier reached an accuracy of 76.9%. Overall, the ANN demonstrated best performance in detecting heart failure, especially with a large number of hidden neurons, highlighting its potential as an effective diagnostic tool.

1. INTRODUCTION

The World Health Organization (WHO) has classified heart disease as one of the leading causes of death worldwide. From the data heart disease observations, approximately 17.9 million lives annually, accounting for roughly 31% of all global deaths (Luo et al., 2024). One critical manifestation of heart disease is heart failure (HF) which is a condition when the heart cannot pump the blood smoothly to reach the body's requirement for oxygen and nutrients. The heart failure happens when the heart becomes weak to sustain their workload, leading to symptoms such as dyspnea (shortness of breath), reduced exercise tolerance, and

^{1*} Corresponding author. E-mail address: norpah020@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.544>

fluid retention. This fluid buildup often results in pulmonary and peripheral edema (Malik et al., 2025). While "heart failure" and "congestive heart failure" are terms often used interchangeably, the latter refers to a specific, more severe form of HF that typically requires urgent medical attention. HF is a chronic and progressive condition that significantly impairs quality of life and increases the risk of mortality. It reflects a weakened heart muscle that struggles to pump sufficient blood throughout the body. As a result, individuals with HF often experience a diminished quality of life and increased healthcare needs. Due to its high prevalence, significant mortality rate, and substantial healthcare costs, heart failure remains a major global health concern. Early and accurate diagnosis is essential to managing the disease effectively (Porumb et al., 2020). Despite medical advances, HF continues to be one of the most costly and deadly health conditions worldwide (Reddy et al., 2021).

HF is influenced by several interconnected factors, including anemia, age, blood pressure, creatine phosphokinase (CPK), and diabetes. Anemia worsens HF by reducing oxygen delivery, thereby increasing cardiac workload and leading to poorer outcomes. It is commonly associated with chronic kidney disease (CKD) and diabetes, creating a complex interplay that heightens cardiovascular risk. Older adults are at a higher risk of developing heart failure due to common risk factors such as hypertension and coronary artery disease. As these conditions become more prevalent with age, the incidence of heart failure is expected to continue rising in the coming decades (Vigen et al., 2012). Hypertension acts both as a cause and consequence of HF; uncontrolled blood pressure can lead to arterial stiffness and elevated pulse pressure, which are linked to anemia. In hypertensive patients, anemia may result from reduced red blood cell deformability and the impact of anti-hypertensive medications on erythropoiesis. Additionally, elevated CPK levels, which indicate muscle damage including myocardial injury, are important markers of ongoing cardiac stress or infarction.

Diagnosing and analyzing heart failure (HF) involves numerous interconnected factors, often making it challenging for physicians to reach timely and accurate decisions. To assist in this complex process, Artificial Neural Networks (ANN) have become valuable tools in medical applications, particularly in HF detection. Inspired by the structure of the human brain, ANN are a type of deep neural network capable of modeling non-linear and complex relationships. These systems are composed of layers of interconnected artificial neurons: input, hidden, and output layers where signals are transmitted, processed, and propagated. A key strength of ANN lies in their ability to handle large and complex datasets. By learning from training data, they continuously improve their accuracy, making them highly effective for classification and clustering tasks (IBM Cloud Education, 2020). While traditional machine learning methods have also been applied to HF prediction, ANN offer superior adaptability and learning capabilities. However, a notable drawback is that the size and structure of ANN architectures are typically determined by trial and error or prior experience, and their performance heavily depends on selecting the right configuration. Adding new features to the system can also affect its accuracy (Maad, 2021).

In industrial contexts, ANN provides cost-effective solutions by reducing material usage (Filippis et al., 2018). In healthcare, their predictive power can support early intervention, potentially lowering the need for expensive surgeries and treatments. To harness these benefits, healthcare systems require advanced data storage, analysis, and data mining techniques to extract meaningful insights from vast datasets (Santos-Pereira et al., 2022). While data mining has already proven successful in fields like marketing and retail, its application in healthcare is still evolving. As big data continues to grow, neural networks are revolutionizing both industry and daily life by advancing artificial intelligence. Given the rising mortality rates from heart failure, ANN are especially suited for big data-driven medical applications, helping physicians predict disease outcomes and manage large volumes of patient data, thereby reducing the burden on clinical experts.

This research focuses on accurately predicting the presence of HF using two distinct algorithms: Scaled Conjugate Gradient (SCG) and Naïve Bayes (NB). The proposed models utilize only 13 attributes (as shown in Table 1). These observed features form the core basis for testing, which generally yields reliable results. SCG is an optimization method employed for training neural networks, offering benefits such as faster

convergence compared to traditional backpropagation and greater efficiency when handling large datasets. On the other hand, Naïve Bayes is a simple, yet effective probabilistic classifier based on Bayes' Theorem, operating under the “naïve” assumption that features are conditionally independent given the class label. Despite this often unrealistic assumption, Naïve Bayes has demonstrated strong performance in medical applications (Girtonia et al., 2019).

Neural Networks are one of the most powerful and widely used machine learning models today. Neural Network models are inspired by the structure of the human brain, capable of capturing complex non-linear relationships in data. The ANN with an appropriate network structure can handle the correlation or dependence between input variables and is known for its ability to learn complex patterns and relationships in data. Furthermore, the Scaled Conjugate Gradient (SCG) method is a neural network training algorithm. Basim et al. (2025) in their research have proved the effectiveness of these methods is illustrated in the context of training artificial neural networks. Experimental results show that the improvement SCG methods achieves competitive performance in terms of convergence rate and accuracy.

Naive Bayes are probabilistic model used in machine learning and data analysis. Naive Bayes assumes that all features are independent of each other, given the class label. This simplifying assumption makes Naive Bayes computationally efficient and easy to implement. Naive Bayes is often used in text classification and spam filtering tasks due to its simplicity and good performance in practice. Naive Bayes algorithms accurately categorize news headlines into different classes. This algorithm is well-established and widely used machine learning algorithm known for its simplicity and effectiveness in text classification tasks (Merve et al., 2024).

Thus, the objective of this research is to identify highly accurate methods for detecting heart failure. The sub objectives are to use structured data to predict the risk of heart failure and to use deep learning to determine the specific risk of heart failure specifically on Neural Network and Naïve Bayes.

The study's findings highlight the critical importance of early detection in heart failure, which can mean the difference between life and death. Recognizing the early signs of heart failure significantly improves our ability to detect threats in time. Medical history remains one of the most reliable indicators for predicting the likelihood of developing heart disease. This study is particularly significant because it helps identify appropriate treatment targets, advancing our goal of improving both the quality and longevity of life for patients with heart failure. Timely detection is essential, as the condition progressively worsens without proper intervention. Therefore, early diagnosis is crucial to ensure effective treatment. Additionally, many existing treatments for heart failure remain under-researched. Health authorities must prioritize heart failure management to reduce its growing economic burden and improve patient outcomes.

The rest of this paper is organized as follows. The second section explains on the data and methodology, consisting of essential steps of Neural Network and Naïve Bayes. In the third section present the results, findings and discussion of the analysis. Finally, the fourth section presents the conclusions and future works.

2. METHODOLOGY

This study employed a quantitative research methodology to develop and evaluate predictive models for heart failure (HF) detection. The approach involved using two distinct classification algorithms: an Artificial Neural Network (ANN) and a Naïve Bayes classifier, to predict a patient's 'death event' during a follow-up period. The following sections detail the experimental setup, from data acquisition and preprocessing to model training and performance evaluation.

2.1 Dataset information

The research utilized secondary data from the Heart Failure Clinical Records Dataset, sourced from the UCI Machine Learning Repository (UCI Machine Learning Repository, 2025). The dataset comprises 13 attributes from the medical records of 299 patients. It contains 12 input features and one target variable. The input features, which include a mix of demographic data, physiological measurements, and clinical observations, are detailed in Table 1 below. This rich feature set allows the models to learn complex relationships between patient characteristics and heart failure outcomes. The target variable, Death event, is a crucial Boolean attribute indicating whether a patient passed away during the follow-up period.

Table 1. Overview of Attributes

Attributes	Description
Age	Age of the patient (years)
Anemia	Decrease of red blood cells or hemoglobin (Boolean)
High blood pressure	If the patient has hypertension (Boolean)
Creatinine phosphokinase (CPK)	Level of the CPK enzyme in the blood (mcg/L)
Diabetes	If the patient has diabetes (Boolean)
Ejection fraction	Percentage of blood leaving the heart at each contraction (percentage)
Platelets	Platelets in the blood (kilo platelets/mL)
Sex	Woman or man (binary)
Serum creatinine	Level of serum creatinine in the blood (mg/dL)
Serum sodium	Level of serum sodium in the blood (mEq/L)
Smoking	If the patient smokes or not (Boolean)
Time	Follow-up period (days)
Death event (target)	If the patient deceased during the follow-up period (Boolean)

Source: UCI Machine Learning Repository (2025)

2.2 Artificial Neural Network (ANN) Implementation

The Artificial Neural Network (ANN), specifically a Multi-Layer Perceptron (MLP), was designed as a two-layer Feed-Forward Neural Network (FFNN). The implementation of the ANN followed a systematic process:

- Step 1: Data Preprocessing and Partitioning - Categorical attributes were converted to numerical values using dummy variables. The 299 samples were then randomly divided into three distinct categories: 70% for training (209 samples), 15% for validation (45 samples), and 15% for testing (45 samples).
- Step 2: Model Architecture - The ANN architecture was defined with 12 input neurons, 20 hidden neurons, and a single output neuron. The network used sigmoid hidden neurons and softmax output neurons.
- Step 3: Model Training - The training process employed the Scaled Conjugate Gradient (SCG) backpropagation algorithm. Training was automatically terminated when the generalization performance ceased to improve, indicated by an increase in the cross-entropy error of the validation samples.
- Step 4: Prediction Generation - The trained model was used to predict outcomes on the unseen testing dataset.
- Step 5: Performance Evaluation - The overall testing results were evaluated using several key metrics, including the confusion matrix, error histogram, and receiver operating characteristic (ROC) curves, to assess the model's accuracy, precision, and recall.

Fig. 1 presents the flowchart of ANN process.

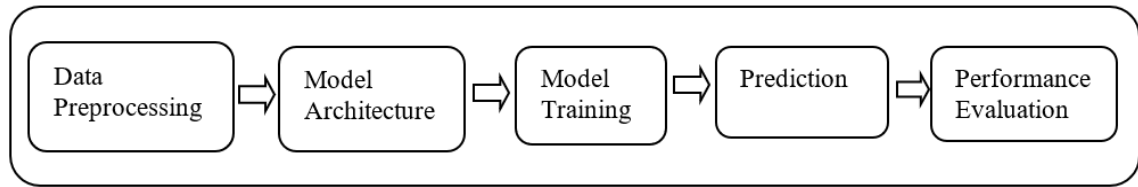


Fig. 1. Flowchart of ANN process

2.3 Naïve Bayes (NB) algorithm implementation

The Naïve Bayes (NB) classifier was chosen as a second model for its probabilistic foundation and computational simplicity. The algorithm is grounded in Bayes' probability theorem and operates under the "naïve" assumption that its features are conditionally independent given the class label. Despite this often-unrealistic assumption, Naïve Bayes has demonstrated strong performance in medical applications. The implementation of the Naïve Bayes algorithm in this study followed a systematic process:

- Step 1: Data Acquisition and Preparation - The dataset was imported into the analytical software. No missing or erroneous values were found.
- Step 2: Data Partitioning - The dataset was divided into distinct training and testing sets, incorporating a 5-fold cross-validation to ensure a robust model evaluation.
- Step 3: Model Selection and Training - The Naïve Bayes classifier was selected and applied to the training dataset, where the algorithm learns the probabilistic relationships within the data.
- Step 4: Prediction Generation - The trained model was then used to predict outcomes on the unseen testing dataset.
- Step 5: Performance Evaluation - Finally, the model's accuracy was assessed by comparing the predicted outcomes with the actual values in both the training and testing sets.

Fig. 2 presents the flowchart of Naïve Bayes process.

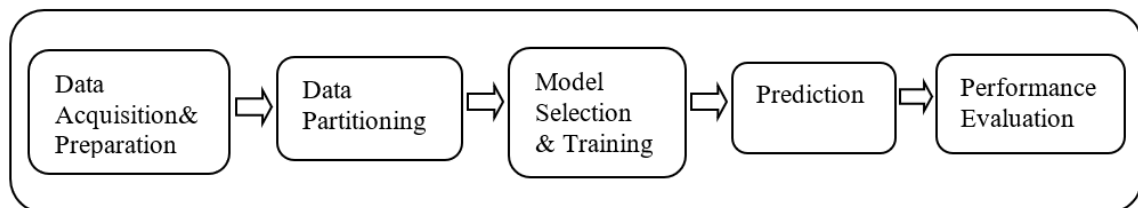


Fig. 2. Flowchart of Naïve Bayes process

2.4 Baseline Method

The Naïve Bayes classifier served as the baseline method for this study. As a probabilistic model with a strong assumption of feature independence, its performance provides a benchmark against which the more complex ANN can be compared. By establishing this baseline, the research can determine if the increased architectural complexity and computational cost of the ANN lead to a statistically significant improvement in predictive accuracy.

2.5 Reproducibility of the proposed work

To ensure the reproducibility of the research, the experimental setup was meticulously documented. The models were implemented using the MATLAB software environment, leveraging its built-in tools such as the Neural Network Pattern Recognition Tool (nprtool) for data assignment and the SCG training function for the ANN model. All random seeds were fixed to a constant value, which guarantees that data splits and model initializations remain identical across multiple runs. This practice eliminates a source of variability, ensuring that the results are consistent and can be replicated by other researchers.

2.6 Performance metrics

The performance of both models was evaluated using several key metrics derived from a confusion matrix. This approach is particularly important in medical applications, where the costs associated with different types of errors vary significantly. A confusion matrix is a table, as shown in Table 2, that summarizes the results of a predictive model by counting the number of correct and incorrect predictions. It quantifies outcomes into four categories:

- True Positive (TP): Positive instances correctly identified.
- True Negative (TN): Negative instances correctly predicted as negative.
- False Positive (FP): Negative instances incorrectly classified as positive.
- False Negative (FN): Positive instances incorrectly classified as negative.

Table 2. Confusion matrix of the predictive model

		Actual Values (Target)	
		0	1
Predicted Values (Output)	0	TN	FN
	1	FP	TP

Source: Adapted from Swaminathan and Tantri (2024).

The four main metrics used to evaluate the model's performance are:

- Accuracy: The overall proportion of correct predictions. While a useful starting point, it can be misleading in cases of class imbalance.
- Precision: The ratio of true positives to all positive predictions. High precision is crucial when a false positive (incorrectly predicting a death event) could lead to unnecessary emotional distress or costly interventions.
- Recall (Sensitivity): The ratio of true positives to all actual positive cases. Recall is arguably the most critical metric in this study, as a false negative (failing to predict a death event) has severe, potentially life-threatening consequences.
- F1-Measure: The harmonic mean of precision and recall. It provides a single metric that balances the trade-off between the two.

The formulas for these metrics are:

$$Recall = \frac{TP}{TP + FN} \quad (1)$$

$$Accuracy = \frac{TP + TN}{TP + FP + TN + FN} \quad (2)$$

$$Precision = \frac{TP}{TP + FP} \quad (3)$$

$$F1 - measure = \frac{2 \times Precision \times Recall}{Precision + Recall} \quad (4)$$

These metrics provide a nuanced understanding of the model's predictive capabilities

3. RESULTS AND DISCUSSION

A total of 209 patient data samples were used in the training of the Artificial Neural Network (ANN), with 45 samples used for validation. The input model has been trained for testing using Scaled Conjugate Gradient (Trainscg) algorithm.

The ANN training tool is shown in Fig. 3 below. Once the tool has ended, the model can open the plot. The ANN architecture shows that the model has a two-layer Feed-Forward Neural Network with sigmoid hidden and *softmax* output neurons. Besides, the training of the Trainscg algorithm shows the model training function updates the weight and bias value according to the Scaled Conjugate Gradient optimization. The ANN training will evaluate the best validation performance in the ANN and choose the cross-entropy to plot the performance.

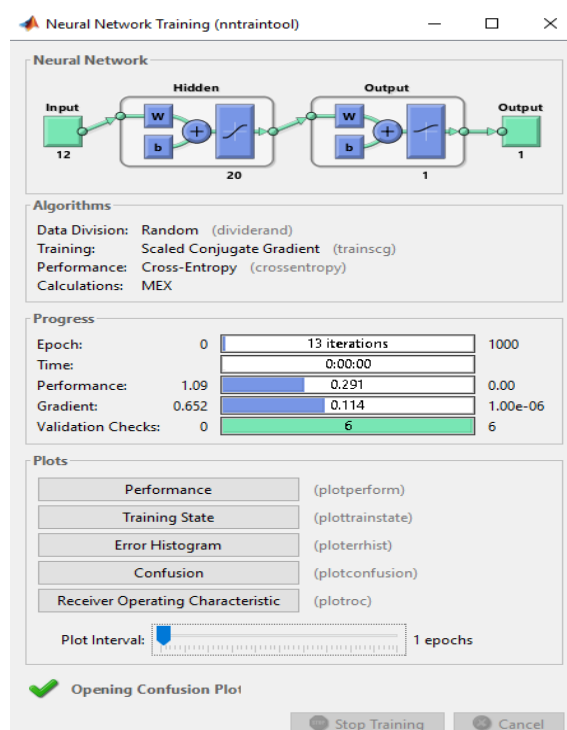


Fig. 3. NN training tool

Fig. 4 shows the trend of the train, validation, and test performance in detail. From this plot, the x-axis represents the epochs, while the y-axis represents Cross Entropy which is the minimization value of our error function. The training curve was seen decreasing with the increase of epoch number till it reached 13 epochs. It shows that the trend has appeared near the best-dotted line and to be specific, it means that the training has been done successfully and converged. The best validation performance is 0.38873 at epoch 7. The training finishes after six consecutive epochs and then increases in validation error in the default configuration. The best performance comes from the epoch with the lowest validation error.

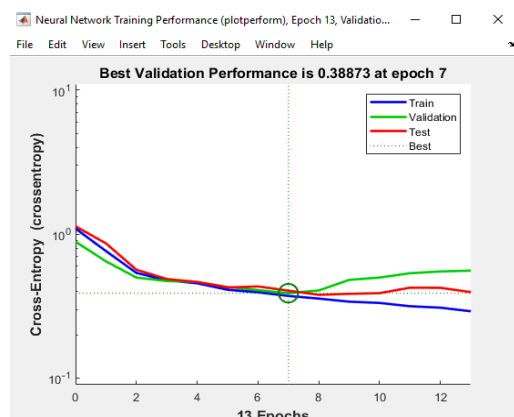


Fig. 4. NN training performance

The ANN training state, shown in Fig. 5, displays two different types of graphs. The epochs are represented on the x-axis, while the gradient value and failed validation check are represented on the y-axis. The network weights are balanced throughout the training phase to reduce the gradient, which depicts the evolution of the 0.11382 minimum gradient. Meanwhile, the second graph shows that the network stops training after six consecutive failures, where the validation check is used to stop artificial neural networks from learning.

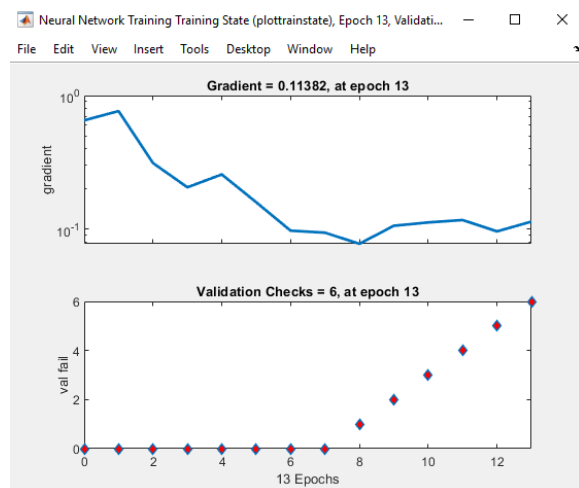


Fig. 5. NN training state

The error in histogram occurred between target and predicted values after training a Feed-Forward Neural Network (FFNN). Based on Fig. 6, the plot shows the error histogram with 20 bins. The bin corresponding to the error is (-0.02027). The height of the bin for training datasets lies below but near 70. The validation and test datasets lie between 70 and 90, and this means the error lies between that range.

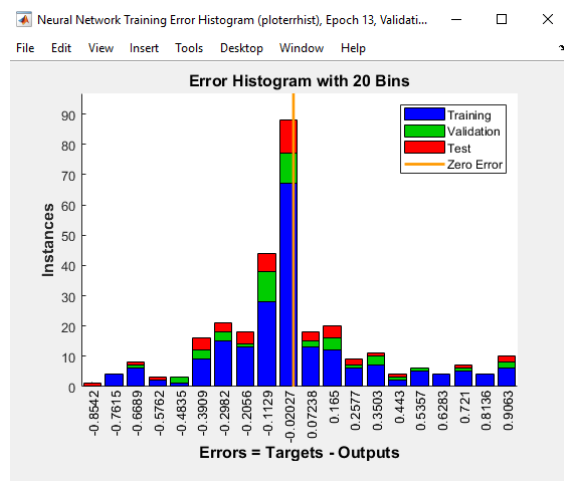


Fig. 6. NN training error histogram

Next is the Receiver Operating Characteristic (ROC) on the ANN and Naïve Bayes. The ROC curve depicts the true-positive (TP) rate versus the false-positive (FP) rate for the learned classifier currently in use. The best feasible prediction approach would result in a point in the upper left corner of the ROC space, with 100% sensitivity and 100% specificity. For example, Fig. 7 shows that the ANN generates the best ROC because it appeared near the coordinate (0,1).

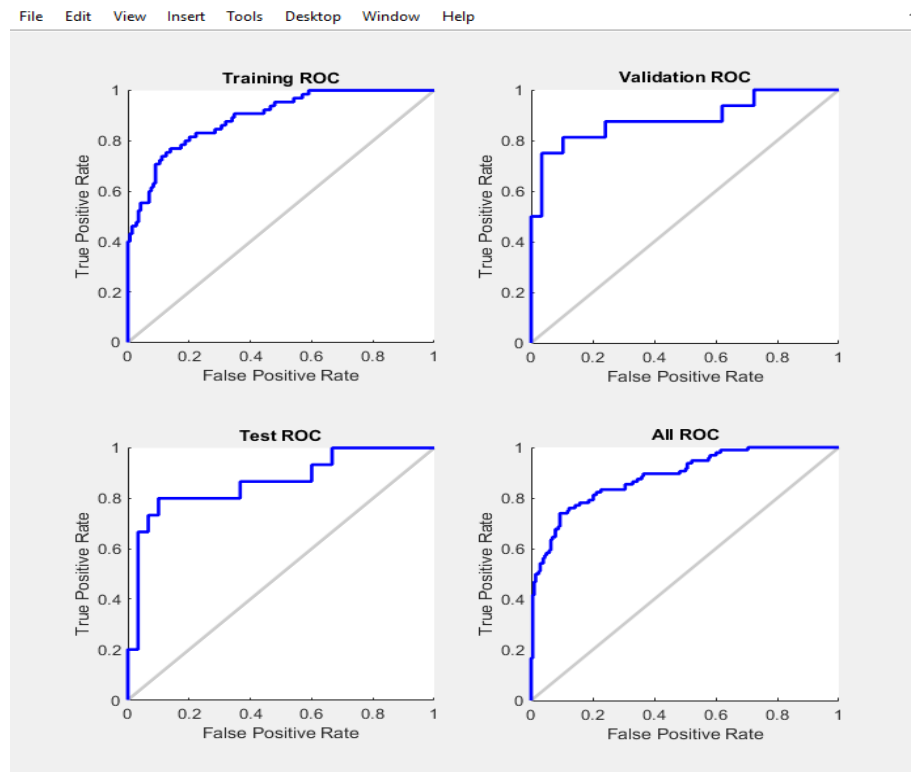


Fig. 7. ROC of ANN

<http://doi.org/10.24191/jcrinn.v10i2.544>

The red marker indicates the performance of the currently selected classifier on the plot. For example, Fig. 8 shows the ROC curve for patients who do not suffer from heart disease, which an FP rate of 0.54 means that the present classifier wrongly allocates 54% of the data to the positive class. On the other hand, with a TP rate of 0.92, the current classifier correctly allocates 92% of the observations to the positive class. It also shows a large Area Under Curve (AUC) of 0.87, indicating better classifier performance.

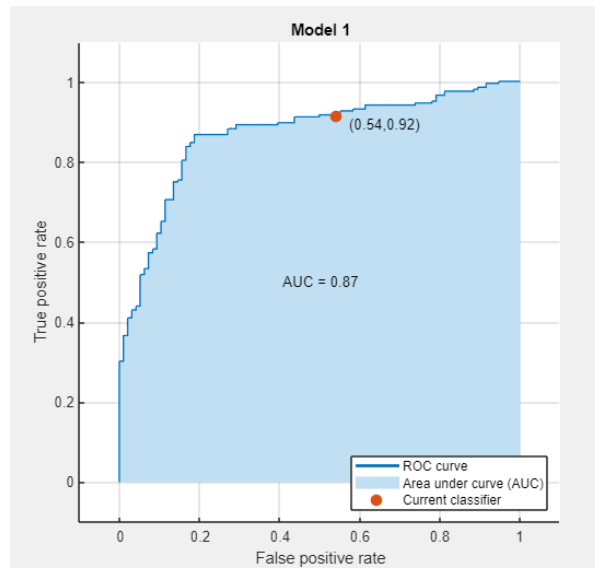


Fig. 8. ROC of Naïve Bayes

Next, the performance criteria will be calculated by employing the Equation (1) - (4) as follows.

ANN:

$$\begin{aligned}
 \text{Recall} &= \frac{27}{27 + 3} \times 100\% &&= 90\% \\
 \text{Accuracy} &= \frac{27 + 12}{27 + 3 + 3 + 12} \times 100\% &&= 86.7\% \\
 \text{Precision} &= \frac{27}{27 + 3} \times 100\% &&= 90\% \\
 \text{F1 - measure} &= \frac{2 \times 0.9 \times 0.9}{0.9 + 0.9} \times 100\% &&= 90\%
 \end{aligned}$$

Naïve Bayes:

$$\begin{aligned}
 \text{Recall} &= \frac{186}{186 + 17} \times 100\% &&= 91.6\% \\
 \text{Accuracy} &= \frac{186}{186 + 44 + 17 + 52} \times 100\% &&= 76.9\% \\
 \text{Precision} &= \frac{186}{186 + 52} \times 100\% &&= 78.2\% \\
 \text{F1 - measure} &= \frac{2 \times 0.916 \times 0.782}{0.916 + 0.782} \times 100\% &&= 84.4\%
 \end{aligned}$$

Based on Table 3 below, ANN shows higher accuracy than Naïve Bayes, 86.7% and 76.9%, respectively. ANN showed high accuracy because it can handle a large amount of dataset provided, yet it can indirectly discover complex non-linear interactions between dependent and independent variables. Even though Naïve Bayes has lower accuracy than ANN, but the result still shows that Naïve Bayes also has a good accuracy range. Naïve Bayes is a simple yet powerful algorithm to be chosen. This can be proven by a previous study by Dangare and Apte (2012) which revealed the comparison results between ANN and Naïve Bayes in predicting heart disease, which is 99.25% and 94.44%, respectively. Similarly, Sajja and Kalluri (2020) have proved the results of the comparison between ANN and Naïve Bayes in their study, which are ANN and Naïve Bayes showed 86.97% and 77.04%, respectively. Overall, this indicates that ANN has better performance for this dataset.

Table 3. Comparison of performance measure of ANN and Naïve Bayes

Algorithm	Sensitivity	Precision	F1 Score	Accuracy
ANN	90.0%	90.0%	90.0%	86.7%
Naïve Bayes	91.6%	78.2%	84.4%	76.9%

4. CONCLUSION

The present study was designed to show that the classification of heart disease is important in our daily lives. This is because it can assist doctors and experts in screening the process of detecting heart failure in the patient. In order to identify highly accurate methods for detecting heart failure, there are two methods used in this study, the Artificial Neural Network (ANN) and the Naïve Bayes. This study found that using ANN with 20 hidden layers gave us good accuracy. The advantage of a ANN is that it can detect complicated non-linear correlations between dependent and independent variables without requiring extensive statistical training. The acquired result demonstrates that the training data created allows the ANN to predict the output for the testing dataset. This study indicates that ANN produced 86.7% testing accuracy while Naïve Bayes showed 76.9% testing accuracy. Overall, ANN showed significant results in the diagnosis of heart disease. This shows us that the results of this study support the idea that deep learning can be used to determine the specific risk of heart failure by using structured datasets.

Implementing a future study exploring the potential of ANN is an inspiring prospect. The research could be broadened to encompass longitudinal studies of patients, thereby enhancing the accuracy of heart disease prediction. This study proposes the application of a Multi-Layer Perceptron Neural Network (MLPNN), which has shown promising results. Its implementation could assist domain experts and individuals in the field in planning for improved diagnosis and providing patients with early diagnosis results. The MLPNN performs admirably, even without retraining, and has the potential to reduce the number of hidden layers while enhancing the number of patient-related attributes. Future studies involving NN should be applied in real-world applications, offering the potential to alleviate the burden on experts and benefit the health community while reducing costs.

5. ACKNOWLEDGEMENTS/FUNDING

The authors would like to acknowledge the Journal of Computing Research and Innovation (JCRINN) for their valuable recommendation on this research.

6. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

7. AUTHORS' CONTRIBUTIONS

Norpah Mahat: Conceptualisation, supervision and writing original draft and investigation; **Norazwana Saidin @ Zubir:** Conceptualisation, methodology and formal analysis; **Jasmani Bidin:** Writing Introduction; **Izleen Ibrahim:** Methodology; **Sharifah Fahriyah:** Results and discussion; **Mohamad Najib Mohamad Fadzil:** Conclusion and future work; **Siti Sarah Raseli:** Writing- review and editing, and validation.

8. REFERENCES

- Basim, A.H., Reem, M., Alaa, L.I. & Hawraz, N.J. (2025). Improved conjugate gradient methods for unconstrained minimization problems and training recurrent neural network. *Engineering Reports*, 7(2), 1-13. <https://doi.org/10.1002/eng2.70019>
- Dangare, C. S. & Apte, S. S. (2012). Improved study of heart disease prediction system using data mining classification techniques. *International Journal of Computer Application*, 47(10), 44-48. <https://doi.org/10.5120/7228-0076>
- Filippis, L. A. C. D., Serio, L. M., Facchini, F., & Mummolo, G. (2018). ANN modelling to optimize manufacturing process. *Advanced Applications for Artificial Neural Networks*, 11, 201-225. <https://doi.org/10.5772/intechopen.71237>
- Giridonia, M., Garg, R., Jeyabashkharan, P., & Minu, M.S. (2019). Cervical cancer prediction using Naïve Bayes Classification. *International Journal of Engineering and Advanced Technology (IJEAT)*, 8(4), 784-787.
- IBM Cloud Education. (2020). What are neural networks?. *IBM Cloud*.
- Luo, Y., Liu J., Zeng, J. & Pan, H. (2024). Global burden of cardiovascular diseases attributed to low physical activity: An analysis of 204 countries and territories between 1990 and 2019. *American Journal of Preventive Cardiology*, 17, 100633. <https://doi.org/10.1016/j.ajpc.2024.100633>
- Malik, A., Brito, D., & Chhabra., L. (2025). Congestive heart failure. *StatPearls*. StatPearls Publishing. <https://www.statpearls.com/point-of-care/search?q=Congestive+Heart+Failure>
- Merve, V., Erkan, V. & Ihsan, O.B. (2024). Performance comparison between Naïve Bayes and machine learning algorithms for news classification. *Bayesian Inference – Recent Trends*. 1-20. <https://doi.org/10.5772/intechopen.1002778>
- Maad, M. M. (2021). Artificial neural network advantage and disadvantages. *Mesopotamian Journal of Big Data*, 2021, 29-31. <https://doi.org/10.58496/MJBD/2021/006>
- Porumb, M., Iadanza, E., Massaro, S., & Pecchia, L. (2020). A convolutional neural network approach to detect congestive heart failure. *Biomedical Signal Processing and Control*, 55, 101597. <https://doi.org/10.1016/j.bspc.2019.101597>
- Reddy, M.K., Helkkula, P., Keerthana, Y.M., Kaitue, K., Minkinen, M., Tolppanen, H., Nieminen, T., & Alku, P. (2021). The automatic detection of heart failure using speech signals. *Computer Speech and Language*, 69, 101205. <https://doi.org/10.1016/j.csl.2021.101205>

- Sajja, T. K. & Kalluri, H. K. (2020). A deep learning method for prediction of cardiovascular disease using Convolutional Neural Network. *International Information and Engineering Technology Association (IIETA)*, 34(5), 601-606. <https://doi.org/10.18280/ria.340510>
- Santos-Pereira, J., Gruenwald, L., & Bernardino, J. (2022). Top data mining tools for the healthcare industry. *Journal of King Saud University - Computer and Information Sciences*, 34(8), 4968-4982. <https://doi.org/10.1016/j.jksuci.2021.06.002>
- Swaminathan, S., & Tantri, B. R. (2024). Confusion Matrix-Based Performance Evaluation Metrics. *African Journal of Biomedical Research*, 27(4S), 4023-4031. <https://doi.org/10.53555/AJBR.v27i4S.4345>
- UCI Machine Learning Repository. (2025). *Dataset*. UC Irvine Machine Learning Repository. <https://archive.ics.uci.edu/datasets>
- Vigen, R., Maddox, T.M. & Allen, L.A. (2012). Aging of the United States Population: Impact on Heart Failure. *Current Heart Fail Reports*, 9(4), 369–374. <https://doi.org/10.1007/s11897-012-0114-8>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Enhancing Photovoltaic Panel Inspection using RGB Image-Based Detection and Image Processing Techniques

Suhaili Beeran Kutty^{1*}, Mohamad Ad-Fadhil Musa², Murizah Kassim³, Puteri Nor Ashikin Megat Yunus⁴

¹Faculty of Electrical Engineering, Universiti Teknologi MARA Johor Branch, Pasir Gudang Campus, 81750 Masai, Johor, Malaysia.

²SHB Malaysia Automotive Appliance Sdn. Bhd, Johor, Malaysia.

^{3,4}Faculty of Electrical Engineering, Universiti Teknologi MARA, 40450 Shah Alam, Selangor, Malaysia.

ARTICLE INFO

Article history:

Received 6 May 2025

Revised 6 August 2025

Accepted 15 August 2025

Published 1 September 2025

Keywords:

Photovoltaic Panel Defects

RGB Image Analysis

Image Processing

K-means Clustering

Canny Edge Detection

DOI:

10.24191/jcrinn.v10i2.528

ABSTRACT

Photovoltaic (PV) panels have become more common in recent years due to their numerous benefits. However, ensuring that PV panels function optimally and reliably is essential for maximizing their efficiency and durability. Identifying and addressing flaws in PV panels is vital to achieving this goal. While thermographic imaging is routinely utilized for defect identification, the potential for using RGB images for this purpose is virtually untapped. This paper intends to investigate using RGB images to recognize PV panel defects, proposing a methodology that integrates image processing techniques such as K-means clustering, Canny edge detection, and grayscale conversion. The results show that defects on PV panels may be successfully discovered by applying K-means clustering and Canny edge detection to RGB images with an accuracy of 90.66%. This study sheds light on improving defect identification practices in the PV industry.

1. INTRODUCTION

Photovoltaic or PV refers to converting light into electricity using semiconducting materials that exhibit the photovoltaic effect. PV technology is commonly used in solar panels to generate electricity from sunlight. PV technology is considered a key component of renewable energy systems due to its ability to harness abundant sunlight and convert it into clean, sustainable electricity without emitting greenhouse gases or other pollutants (Obaideen et al., 2021). PV systems are widely used worldwide to help reduce reliance on fossil fuels and mitigate climate change by providing a source of clean, renewable energy (Shahsavari & Akbari, 2018). The study by Han et al. (2025) confirmed that utilizing renewable energy enhances energy resilience, whereas dependence on fossil fuels undermines long-term sustainability. The use of solar energy is increasing, including in Malaysia, and its importance needs to be widely promoted to consumers (Amali et al., 2024), as well as how it is implemented. PV panels comprise interconnected

^{1*} Corresponding author. E-mail address: suhaili88@uitm.edu.my
10.24191/jcrinn.v10i2.528

photovoltaic cells, typically made from silicon and other materials. These panels are commonly installed on rooftops, in solar farms, or integrated into various structures to generate renewable electricity for residential, commercial, and industrial applications.

PV technology also does not escape experiencing problems that need to be monitored and maintained, including issues faced by PV panels. Defects on PV panels refer to any abnormalities or faults on panels that can impair their performance, efficiency, or reliability (Osmani et al., 2023). These defects can manifest in various forms and may result from manufacturing flaws, installation issues, environmental factors, or degradation over time. Examples of defects are cracks, hotspots, delamination, soiling, and degradation (Afifah et al., 2021). Detecting and addressing PV defects is crucial for maintaining solar energy systems' optimal operation and longevity.

One of the simplest and most cost-effective methods for detecting defects is visual inspection, which involves manual examination for visible signs of damage like cracks, delamination, or corrosion. However, this method may overlook internal defects or those in hard-to-reach areas. Electroluminescence (EL) imaging, a non-destructive technique, captures images of PV panels under electrical excitation, revealing internal defects such as microcracks, hotspots, and cell shunts (Al-Waisy et al., 2022). Infrared (IR) thermography offers another non-destructive approach by capturing thermal images of panel surfaces to detect temperature variations indicative of localized defects like cell cracks or interconnect failures (Wang et al., 2022a). Electrical characterization techniques involve analyzing electrical parameters such as current-voltage curves and series resistance to identify deviations that may indicate the presence of defects affecting panel performance. Moreover, machine learning (ML) algorithms have emerged as a powerful tool for defect detection in PV panels, automating the identification and classification of defects with high accuracy based on datasets of defect images or electrical characteristics.

Artificial intelligence (AI) has become widely used across various fields, including image detection. However, traditional image processing remains essential for studying the features and characteristics of an image. The study by Et-taleby et al. (2025) used image processing to identify the pixel intensity levels based on the brightness of thermal images, which were then classified into whether the PV panels were damaged. For our work, we report the capability of using image processing techniques without AI assistance to identify whether a panel is defective. Subsequently, classification experiments using AI alone and combining both techniques were conducted. Defect detection using RGB images is somewhat limited, whereas these images can be captured using high-definition cameras built into drones. RGB images are more convenient for computational processes than hyperspectral images (Purwadi et al., 2023). This research examines the feasibility and usefulness of employing RGB image processing techniques to detect defects in PV panels.

Work in (Patel et al., 2020) proposed a method to detect PV defects using image processing based on RGB images. Their work uses thresholding, morphological erosion, and edge detection to detect the defect PV. For edge detection, the Kirsch operator was used to detect the edges and boundaries of the image. RGB images were turned into HSV, grayscale, and binary images for processing. Image processing has also been used by Qasem et al. (2016), where the PV images are filtered, and the pixels' lightness is differentiated to determine the dusty level that appears on the PV.

Next, Salamanca et al. (2017) proposed to detect solar panels after the image was captured by a standard camera without light restriction. In their study, they used computer vision and image processing techniques to identify if there were solar panels in the images. Guan et al. (2023) proposed a method based on the gray-level co-occurrence matrix that allows the defect to be detected based on the images captured by unmanned aerial vehicles to enhance the PV defect detection process.

Two types of images are used to detect the defect: IR and RGB in the work proposed by Kuo et al. (2023). IR imaging was used to detect PV module thermal defects, and RGB was used to detect defects on

the panel's surface. The study used Otsu's method to determine the threshold value and used the Laplacian operator to detect the edge in the image.

Although several studies have applied advanced methods for PV defect detection, these approaches often require expensive equipment, controlled environments, or complex training datasets. Furthermore, RGB image-based methods remain underexplored, particularly for low-cost and real-time applications. This study utilizes RGB images of PV panels to determine whether a panel is defective or not after applying image processing techniques. After this introduction section, this paper will discuss the methodology used in this paper, followed by results and discussion. The final section is a conclusion.

2. METHODOLOGY

In this study, the detection of defects in the RGB image was accomplished by a series of stages that included image pre-processing, histogram analysis, and image segmentation techniques that included two methods: classical segmentation and K-means clustering, edge detection, and panel classification, as shown in Fig. 1.

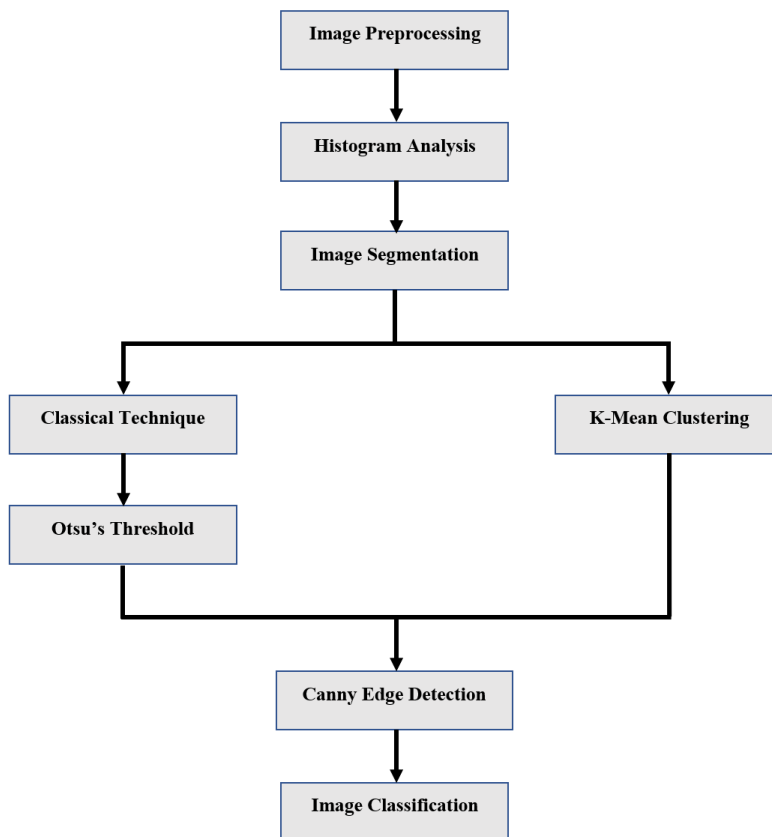


Fig. 1. Flowchart of the process to classify the images into defect and clean PV panels

The selection of Otsu's thresholding, Canny edge detection, and K-means clustering in this study is driven by their proven effectiveness in various image processing applications, particularly in segmentation and edge detection tasks. Otsu's thresholding is a widely used global binarization method that automatically determines the optimal threshold to distinguish foreground from background in grayscale images. Canny edge detection is chosen due to its robustness in detecting sharp intensity changes, which is essential for outlining defect boundaries on PV panels. Meanwhile, K-means clustering is a simple yet powerful unsupervised learning algorithm that segments image regions based on color or intensity similarities, enabling the separation of defective areas from clean regions. These methods are also computationally efficient and adaptable to different image types.

2.1 Image pre-processing

The image data used in this study are publicly available and were obtained from Kaggle. The dataset consists of RGB images of photovoltaic (PV) panels with varying conditions, including clean and defective panels. These images were used to evaluate the effectiveness of image processing techniques for defect detection.

At this preliminary step, the image will go through image cropping and resizing. In this case, the raw image of a panel needs to be cropped and resized to eliminate the background image, leaving only the image of the panel itself. The initial stage of image preprocessing, which involves cropping and resizing images, is designed to optimize the image for further processing. It aims to increase image quality, clarity, and accuracy (Wang et al., 2022b; Zyout & Oatawneh, 2020). During the image preprocessing stage, the background was removed through a cropping process, then the image dimensions were resized from 136×160 pixels to 132×156 pixels.

2.2 Histogram analysis

Image histogram analysis is important in image processing techniques for detecting defects in PV panels (Ali et al., 2020; Prabhakaran et al., 2023). An image's histogram represents the distribution of pixel intensities across the image. Several key insights can be derived from analyzing the histogram, which is helpful for flaw detection in PV panels. The dynamic range of pixel intensities in the image is shown through histogram analysis. Adjusting the contrast based on the histogram might make defects more visible, making them easier to notice. This modification entails expanding or equalizing the histogram to use the entire intensity range of the image. Fig. 2 shows the results of histogram analysis with examples of defective and clean panel images. Regarding light intensity range and the highest pixel, the clean panel image and the defective panel image provide distinct results.

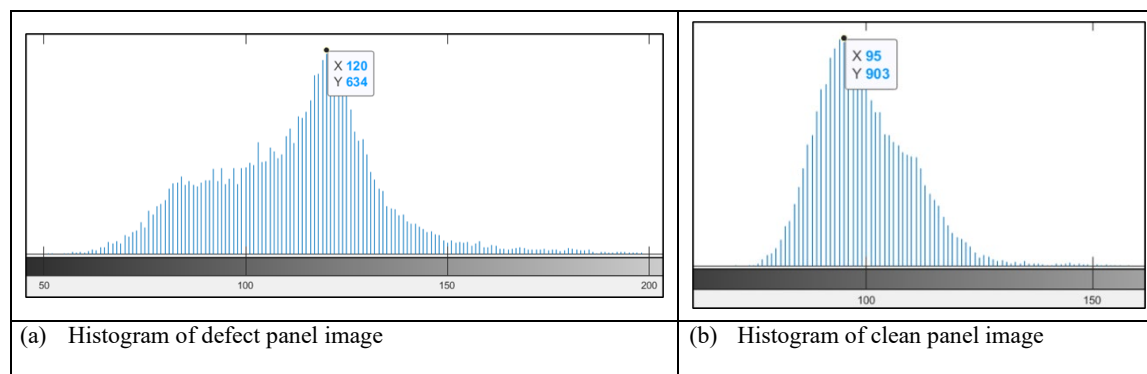


Fig. 2. Histogram analysis of defect and clean panel images

2.3 Image segmentation

Image segmentation refers to splitting the PV panel image into meaningful and distinct parts or segments, depending on specific features, when detecting defects on PV panels. It is an important stage in defect identification since it separates the faults from the background and other normal sections of the panel. In this section, the methods that will be tested on the images are the classical technique (grayscale technique) and k-means clustering.

2.3.1 Otsu's threshold

After the preprocessing stage, the RGB image will be transformed into a grayscale image. The transformed grayscale image simplifies the analysis by lowering the image to a single channel intensity value. The image is then segmented into regions of interest using Otsu's thresholding technique, which is determined by maximizing the inter-class variance. This aids in recognizing faults from normal and background areas (Patel et al., 2020; Espinosa et al., 2020). After thresholding, the Canny edge detection approach is implemented to recognize the edges of the defect image. The study by Fadzli et al. (2024) concluded that the Canny operator emerged as the most effective and comprehensive edge detection technique. Canny edge detection identifies sharp intensity shifts typically associated with defect boundaries. Potential fault areas are identified by recognizing these edges, providing crucial information for the next stage. The grayscale threshold obtained from Otsu's approach is stacked with the edges obtained by Canny edge detection to improve the visualization of the reported faults. This overlay combines binary defect segmentation with edges to produce a more realistic picture of fault boundaries. Fig. 3 depicts the process.

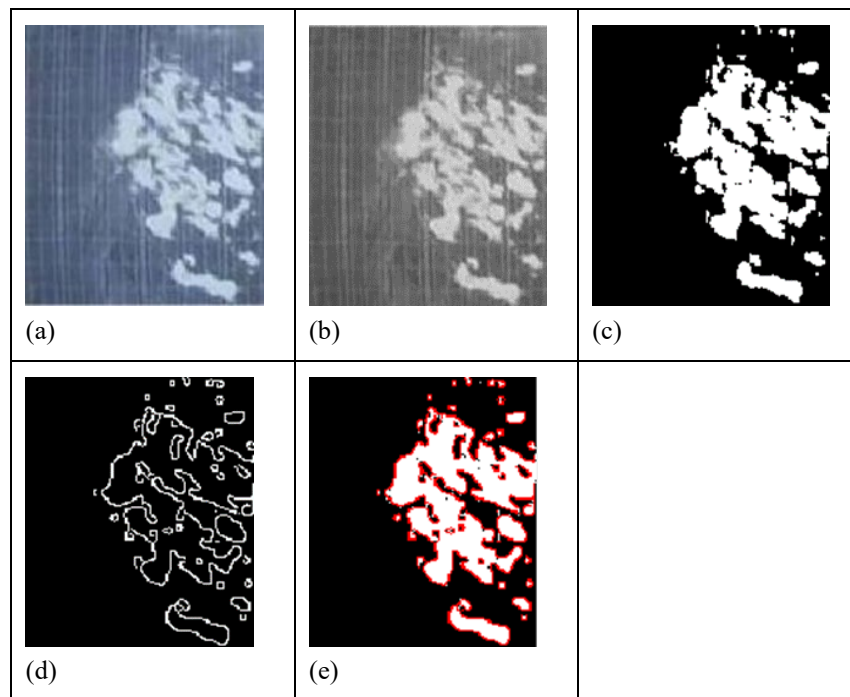


Fig. 3. Image segmentation using Otsu's threshold method. (a) The original image (b) The grayscale image (c) After applying Otsu's threshold method (d) Edge detection using Canny operator (e) Overlay image of (c) and (d)

2.3.2 K-mean clustering

Clustering divides a given set of data points into K unique groups based on similarity or closeness. K-means clustering is used in image processing to group pixels in an image into K clusters, allowing the identification of various regions or objects within the image. In this project, the value of cluster K is set to 2, as the approach of this clustering technique is to separate the image into two clusters, potentially separating defective areas from the rest of the panel. Following the K-means clustering process, the Canny edge detection technique determines the image's edges. Edge detection detects areas with significant intensity shifts, frequently indicative of probable defective areas. The K-means clustered image is overlaid with the edges acquired from the Canny edge detection to improve the visualization of the discovered defects. This approach aids in separating defects from the background and graphically shows probable defect borders, making future defect analysis and characterization easier. Fig. 4 depicts the entire process.

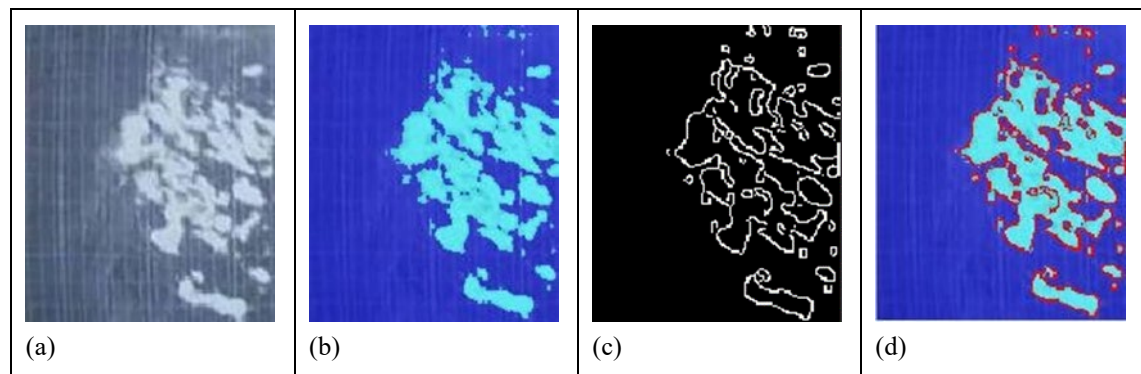


Fig. 4. Image segmentation using the K-means method. (a) The original image (b) K-mean clustering is applied to the image. (c) Edge detection (d) Overlay image of (b) and (c)

2.3.3 Image classification

The number of edge pixels is calculated after applying the classical technique, which includes grayscale conversion, Otsu's method thresholding, and Canny edge detection, as well as the K-means clustering technique, which includes K-means clustering, Canny edge detection, and overlaying the clustered image with the detected edges. Calculating the number of edge pixels offers a quantitative measure of the detected edges, allowing the severity or extent of faults in the PV panel image to be assessed. By comparing the number of edge pixels between a clean PV panel image and a defective PV panel image, it is possible to evaluate the panel as clean or defective by counting the number of edge pixels within the segmented defect regions obtained from either the classical technique or the K-means clustering technique. The number of edge pixels is obtained by traversing the image and counting the pixels detected as edges during the Canny edge detection step when each technique was overlaid. The count is computed within defective areas segmented using either a classical grayscale technique or K-means clustering approaches.

3. RESULTS AND DISCUSSION

In this research, Otsu's threshold and K-means clustering techniques were employed to effectively visualize and analyze defects observed in PV panels using RGB images. By extracting the number of edge pixels from these images, valuable insights into the condition of the panels were obtained. A systematic

comparison of the edge pixel counts between clean and defective panel images established a robust threshold for distinguishing between the two states. This threshold functions as a classification criterion, facilitating the automated recognition and categorization of newly captured images. The findings contribute to the ongoing analysis of PV panels and informed decision-making regarding maintenance and repair strategies.

A dataset of RGB images featuring both clean and defective panels was employed to achieve this purpose. The inquiry began with an initial analysis of the dataset, with a focus on the histogram analysis and the number of edge pixels in each K-means clustering image, as it shows a clearer visual of separating the defective area and the normal area. During the project's preliminary phase, the defect panel image had a low maximum pixel value but a wide light intensity spectrum. In contrast, the clean panel had a high maximum pixel but a small light intensity spectrum. Then, a separate test was done on a dataset of 15 images containing both clean and defective panels before running the images in the graphical user interface developed using MATLAB, as shown in Fig. 5. The objective was to measure the number of edge pixels in each image and provide a preliminary observation. The dataset was purposefully chosen to include many clean and defective panel conditions.

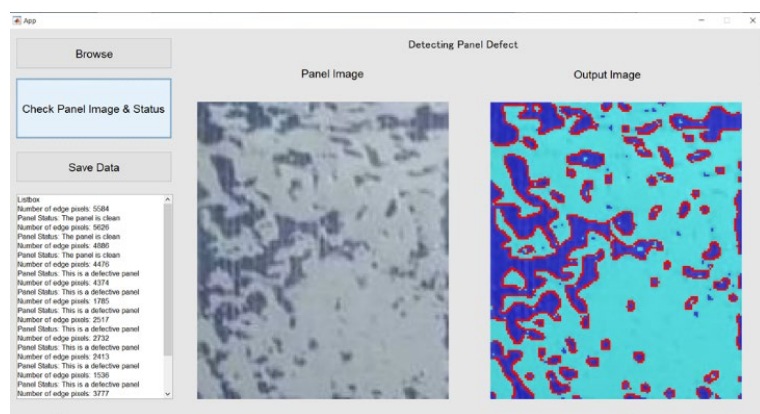


Fig. 5. Graphical User Interface designed using MATLAB for the PV panel detection

After analysing the 15 images, a distinct pattern emerged. The number of edge pixels in clean panel images was consistently 4500 or higher, but the number of edge pixels in defective panel images was 4500 or below. This preliminary discovery gave critical information about the expected distribution of edge pixel counts between clean and damaged panels. The detection and classification were then evaluated on the 610 sample images. The status of defect detection in each image was compared to the 610 sample images labeled as 'actual status' with 'predicted status' as in Table 1 to evaluate the accuracy, precision, and sensitivity of the algorithm used for defect identification in PV plant panels based on RGB images. Positive is classified as a clean panel image, while negative is classified as a defective one.

The data of the total count of true positive (TP), true negative (TN), false positive (FP), and false negative (FN) is in Table 2. The data analysis from Table 3 gives important predictive performance and the actual condition of the panel images. The algorithm achieves an accuracy of 90.66%, indicating that 90.66% of all panel images were correctly classified. This high accuracy demonstrates the algorithm's ability to differentiate between defective and clean panels using RGB image data. The true positive rate or recall is 95.45%, reflecting the algorithm's ability to identify 95.45% of defective panels correctly.

Table 1. Comparison of predicted and actual status in PV panel defect detection

	Predicted Status	Actual Status
True Positive (TP)	Clean	Clean
True Negative (TN)	Defect	Defect
False Positive (FP)	Clean	Defect
False Negative (FN)	Defect	Clean

A high recall value indicates that the system effectively minimizes missed detections of defective panels, thereby enhancing reliability. Similarly, the true negative rate representing the specificity is 90.28%, signifying that 90.28% of clean panels were correctly identified. This high specificity suggests a low false positive rate, reducing the likelihood of misclassifying clean panels as defective. The positive precision is 43.30%, representing the proportion of correctly classified defective panels out of all panels predicted as defective. The relatively lower precision suggests that a significant number of clean panels were misclassified as defective. This could be attributed to an imbalance in the dataset, where defective panels constitute a smaller proportion of the total samples, leading to a higher false positive rate. Conversely, the negative precision is 99.61%, indicating that 99.61% of the panels predicted as clean were indeed clean. This high precision highlights the algorithm's effectiveness in correctly identifying clean panels while minimizing false positives.

Table 2. Total counts of classification results

Total TP	42
Total TN	511
Total FP	55
Total FN	2

Table 3. Evaluation metrics for classification performance

Accuracy	90.66%
True Positive Rate	95.45%
True Negative Rate	90.28%
Positive Precision Value	43.30%
Negative Precision Value	99.61%

4. CONCLUSION

The focus of this research was to detect defects in PV panels using image processing techniques based on RGB imagery. The algorithm separated the RGB images into distinct regions using image processing, allowing the identification of probable defects. Furthermore, the number of edge pixels was used to characterize the existence of defects. The significance of this study lies in its potential to revolutionize defect detection practices in PV plants and significantly impact the solar energy industry. This study makes

several important advances by utilizing RGB image processing techniques for identifying defects in PV panels. For future development planning, the use of AI can be a valuable addition to improving the defect detection system, as proposed by (Othman et al., 2024). By introducing AI into the defect detection system, it is possible to achieve even higher accuracy rates and improve the overall performance of the algorithm. It can also address the dataset imbalance issue identified in this study. This may include applying data augmentation, synthetic image generation, or resampling methods to improve model robustness when classifying underrepresented defective panel images.

5. ACKNOWLEDGEMENTS/FUNDING

The authors would like to express their sincere appreciation to the Faculty of Electrical Engineering and the Research Management Centre (RMC), Universiti Teknologi MARA, for their support. This research was funded by Universiti Teknologi MARA under Grant No. 600-RMC 5/3/GPM (022/2022).

6. CONFLICT OF INTEREST STATEMENT

The authors declare that there is no conflict of interest regarding the publication of the paper.

7. AUTHORS' CONTRIBUTIONS

Suhaili Beeran Kutty: Conceptualisation, methodology, writing, analysis, formal analysis, and project administration; **Mohamad Ad-Fadhil Musa:** Data processing, analysis, and formal analysis; **Murizah Kassim:** Review, editing, and proofreading; **Puteri Nor Ashikin Megat Yunus:** Review, editing, and proofreading.

8. REFERENCES

- Affifah, A. N. N., Indrabayu, Suyuti, A., & Syafaruddin. (2021). A review on image processing techniques for damage detection on photovoltaic panels. *ICIC Express Letters*, 15(7), 779–790. <http://dx.doi.org/10.24507/icicel.15.07.779>
- Ali, M. U., Khan, H. F., Masud, M., Kallu, K. D., & Zafar, A. (2020). A machine learning framework to identify the hotspot in photovoltaic module using infrared thermography. *Solar Energy*, 208, 643–651. <https://doi.org/10.1016/j.solener.2020.08.027>
- Al-Waisy, A. S., Ibrahim, D. A., Zebari, D. A., Hammadi, S., Mohammed, H., Mohammed, M. A., & Damaševičius, R. (2022). Identifying defective solar cells in electroluminescence images using deep feature representations. *PeerJ Computer Science*, 8, e992. <https://doi.org/10.7717/peerj-cs.992>
- Amali, N. A. K., Yahya, N., Taib, M. A. M., Tan, G. J., & Yusof, E. M. M. (2024). Developing a web-based visualization tool for solar energy awareness in Malaysia. *Journal of Computing Research and Innovation*, 9(2), 396–417. <https://doi.org/10.24191/jcrinn.v9i2.466>
- Espinosa, A. R., Bressan, M., & Giraldo, L. F. (2020). Failure signature classification in solar photovoltaic plants using RGB images and convolutional neural networks. *Renewable Energy*, 162, 249–256. <https://doi.org/10.1016/j.renene.2020.07.154>

- Et-taleby, A., Chaibi, Y., Ayadi, N., Elkari, B., Benslimane, M., & Chalh, Z. (2025). Enhancing fault detection and classification in photovoltaic systems based on a hybrid approach using fuzzy logic algorithm and thermal image processing. *Scientific African*, 28, e02684. <https://doi.org/10.1016/j.sciaf.2025.e02684>
- Fadzli, W. M. R. W., Dak, A. Y., & Razak, T. R. (2024). A survey on various edge detection techniques in image processing and applied disease detection. *Journal of Computing Research and Innovation*, 9(2), 23-32. <https://doi.org/10.24191/jcrinn.v9i2.415>
- Guan, Y., Wu, G., Huang, W., Yan, J., Wang, L., & Yi, Y. (2023). Gray level co-occurrence matrix-based defect detection method for photovoltaic power plant panels. In *Proceedings of 2023 International Conference on Computers, Information Processing and Advanced Education CIPAE 2023* (pp. 703-707). IEEE. <https://doi.org/10.1109/CIPAE60493.2023.00136>
- Han, S., Kamaruddin, B. H. B., Shi, X., & Zhu, J. (2025). Achieving energy resilience: Studying renewable and fossil fuel energy generation drivers and COPE-28 pathways of China. *Energy Strategy Reviews*, 58, 101669. <https://doi.org/10.1016/j.esr.2025.101669>
- Kuo, C. F. J., Chen, S. H., & Huang, C. Y. (2023). Automatic detection, classification and localization of defects in large photovoltaic plants using unmanned aerial vehicles (UAV) based infrared (IR) and RGB imaging. *Energy Conversion and Management*, 276, 116495. <https://doi.org/10.1016/j.enconman.2022.116495>
- Obaideen, K., AlMallahi, M. N., Alami, A. H., Ramadan, M., Abdelkareem, M. A., Shehata, N., & Olabi, A. G. (2021). On the contribution of solar energy to sustainable developments goals: Case study on Mohammed bin Rashid Al Maktoum Solar Park. *International Journal of Thermofluids*, 12, 100123. <https://doi.org/10.1016/j.ijft.2021.100123>
- Osmani, K., Haddad, A., Lemenand, T., Castanier, B., Alkhedher, M., & Ramadan, M. (2023). A critical review of PV systems' faults with the relevant detection methods. *Energy Nexus*, 12, 100257. <https://doi.org/10.1016/j.nexus.2023.100257>
- Othman, N. S., Ramli, S., Kamarudin, N. D., Mohamad, A. U., & Ong, A. T. (2024). Classification of defect photovoltaic panel images using matrox imaging library for machine vision application. *International Journal on Informatics Visualization*, 8(3-2), 1528–1535. <https://dx.doi.org/10.62527/joiv.8.3-2.2182>
- Patel, A. V., McLauchlan, L., & Mehrubeoglu, M. (2020, December). Defect detection in PV arrays using image processing. In *2020 International Conference on Computational Science and Computational Intelligence (CSCI)* (pp. 1653–1657). IEEE. <https://doi.org/10.1109/CSCI51800.2020.00304>
- Purwadi, Abu, N. A., Mohd, O. B., & Kusuma, B. A. (2023). Mixed pixel classification on hyperspectral image using imbalanced learning and hyperparameter tuning methods. *International Journal on Informatics Visualization*, 7(3), 910-919. <https://doi.org/10.30630/joiv.7.3.1758>
- Prabhakaran, S., Uthra, R. A., & Preetharoselyn, J. (2023). Deep learning-based model for defect detection and localization on photovoltaic panels. *Computer Systems Science and Engineering*, 44(3), 2683–2700. <https://doi.org/10.32604/csse.2023.028898>
- Qasem, H., Mnatsakanyan, A., & Banda, P. (2016, June). Assessing dust on PV modules using image processing techniques. In *2016 IEEE 43rd Photovoltaic Specialists Conference (PVSC)* (pp. 2066–2070). IEEE. <https://doi.org/10.1109/PVSC.2016.7749993>
- Salamanca, S., Merchán, P., & García, I. (2017, July). On the detection of solar panels by image processing techniques. In *2017 25th Mediterranean Conference on Control and Automation (MED)* (pp. 478-483). IEEE. <https://doi.org/10.1109/MED.2017.7984163>

- Shahsavari, A., & Akbari, M. (2018). Potential of solar energy in developing countries for reducing energy-related emissions. *Renewable and Sustainable Energy Reviews*, 90, 275–291. <https://doi.org/10.1016/j.rser.2018.03.065>
- Wang, X., Yang, W., Qin, B., Wei, K., Ma, Y., & Zhang, D. (2022a). Intelligent monitoring of photovoltaic panels based on infrared detection. *Energy Reports*, 8, 5005–5015. <https://doi.org/10.1016/j.egyr.2022.03.173>
- Wang, Q., Paynabar, K., & Pacella, M. (2022b). Online automatic anomaly detection for photovoltaic systems using thermography imaging and low rank matrix decomposition. *Journal of Quality Technology*, 54(5), 503–516. <https://doi.org/10.1080/00224065.2021.1948372>
- Zyout, I., & Oatawneh, A. (2020, February). Detection of PV solar panel surface defects using transfer learning of the deep convolutional neural networks. In *2020 Advances in Science and Engineering Technology International Conferences (ASET)* (pp. 1–4). IEEE. <https://doi.org/10.1109/aset48392.2020.9118384>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Bibliometric Analysis of Research on Firth Penalized Logistic Regression in Addressing Complete Separation

Nurul Husna Jamian^{1*}, Ahmad Zia Ul-Saufie², Mohammad Nasir Abdullah³

^{1,3} Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Perak Branch, Tapah Campus, 35400 Tapah Road, Perak, Malaysia.

² Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA (UiTM), 40450 Shah Alam, Malaysia.

ARTICLE INFO

Article history:

Received 26 May 2025

Revised 18 August 2025

Accepted 19 August 2025

Online first

Published 1 September 2025

Keywords:

Firth Penalized Logistic Regression
Complete Separation
Logistic Regression
Bibliometric Analysis

DOI:

10.24191/jcrinn.v9i2.535

ABSTRACT

Complete separation in logistic regression leads to infinite estimates which prevents reliable inference. Firth's penalized likelihood method has emerged as a widely accepted and reliable solution that provides finite and more stable estimates. Despite its growing relevance, a thorough understanding of the global research on this topic remains limited. This study conducts a bibliometric analysis of trends related to complete separation in logistic regression using Firth penalized regression. Bibliographic data were retrieved from the Scopus database and analysed using Microsoft Excel and VOSviewer software. After applying inclusion criteria, nine journal articles published between 2012 and 2024 were identified through a structured search conducted on February 22, 2025. The findings reveal a small but growing body of literature, reflecting the emerging status of research on complete separation in logistic regression using Firth penalized regression. The results show an upward trend in publications, particularly from 2019 onward with the United States and Malaysia identified as the most productive countries. Influential articles contributed to methodological development and applications in health and transportation research. Keyword co-occurrence analysis identified thematic clusters in human studies, statistical modelling, and estimation techniques. These findings provide an overview of publication trends, collaboration networks, and research gaps which could support future methodological and multidisciplinary integration of Firth penalized regression.

1. INTRODUCTION

Logistic regression is a widely used statistical method for modelling binary outcome variables, particularly in fields such as medicine, epidemiology, social sciences, and genetics (Gilbert, 2022; Hess & Hess, 2019). It is commonly applied to predict outcomes such as disease occurrence, patient readmission, or treatment efficacy (Gilbert, 2022). The technique has shown impressive results in medical research, aiding in the identification of risk factors, assessment of disease probabilities, and support for clinical decision-making (Olowe et al., 2024). Logistic regression models apply concepts such as odds ratios and logit transformation to examine relationships between predictors and outcomes (Nathanson & Higgins, 2008).

However, logistic regression models can encounter serious estimation issues when the data exhibit a phenomenon known as complete separation. Complete separation arises when a predictor or combination of predictors perfectly discriminates between outcome categories, causing the maximum likelihood estimates (MLE) of regression coefficients to diverge towards infinity (Allison, 2008). This problem compromises the reliability of model interpretation, standard errors, and inference, rendering the traditional logistic regression model unsuitable for such datasets (Botes & Fletcher, 2014).

^{1*} Corresponding author. E-mail address: nurul872@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.535>

To address this challenge, Firth (1993) proposed a penalized likelihood approach that applies a bias-reducing modification to the score function. This method known as Firth's penalized logistic regression that has been demonstrated to effectively resolve issues arising from complete separation without the need to remove covariates (Clark et al., 2023). Firth's approach quickly gained traction due to its ability to produce finite, stable, and more reliable parameter estimates, especially in small sample sizes or sparse data conditions. These characteristics make it particularly useful in fields such as genetics, rare disease research, and case-control study designs (Suhas et al., 2023; D'Angelo & Ran, 2024).

Given its rising prominence and methodological importance, a bibliometric analysis is both timely and warranted. Bibliometric analysis is a quantitative method used to assess academic literature within a specific field, providing a comprehensive overview of publication trends, influential works, author and institutional productivity, collaboration networks, and thematic developments over time (Kumar, 2025). Prior studies have highlighted its methodological strengths, little is known about its bibliometric landscape that is, where, how, and by whom the method has been applied and studied. By examining publication patterns, citation structures, and keyword trends, researchers can identify research gaps, emerging themes, and identify potential areas for future exploration.

The objectives of this study are to conduct a bibliometric analysis of global research related to complete separation in logistic regression using Firth penalized regression. Specifically, the study aims to examine publication trends over time to understand the growth and evolution of interest in this topic. In addition, it intends to identify the most influential articles, authors, and journals. Through this approach, the study seeks to highlight research patterns, shifts, and contributions within the literature, providing a comprehensive understanding of the development and impact of Firth penalized regression in addressing complete separation in logistic regression models.

2. LITERATURE REVIEW

Complete separation in logistic regression occurs when one or more independent variables perfectly predict the binary outcome, leading to infinite or zero maximum likelihood estimates (Botes & Fletcher, 2014; Mansournia et al., 2017). This issue presents a significant challenge, particularly prevalent in small samples or datasets with rare events (Mansournia et al., 2017). Traditional maximum likelihood estimation (MLE) often fails under these conditions, leading to infinite parameter estimates. Firth's penalized likelihood method has been widely adopted to address this issue, offering finite and more reliable estimates (Alam et al., 2022).

Firth (1993) laid the foundational work by introducing a general bias-reduction method for MLEs through penalized likelihood. Although originally proposed for exponential family models, its application in logistic regression proved particularly impactful. The technique modifies the score function to counteract the first-order bias of MLE, thereby producing finite and less biased estimates. Subsequent research by Heinze and Schemper (2002) adapted Firth's approach specifically to logistic regression, demonstrating its superiority over traditional methods when complete or quasi-complete separation occurs. Their simulations revealed that Firth's method not only yields finite estimates but also improved coverage probabilities for confidence intervals, particularly in small-sample settings.

Puhr et al. (2017) evaluated Firth's logistic regression under rare event conditions and found that although it effectively reduced bias in parameter estimates, it tended to skew predicted probabilities toward 0.5 in imbalanced datasets. To address it, they proposed post-hoc and augmentation-based corrections to improve predictive accuracy without undermining bias reduction. Walker and Smith (2019) found that sparseness in binary predictors introduces substantial bias with small sample sizes, which Firth's procedure can effectively correct. Karabon (2020) emphasized Firth's method as a solution for rare events and complete separation, demonstrating its advantages over other approaches such as Fisher's Exact test and Exact logistic regression. Stolte et al. (2024) provided a broader review of bias-reduction methods, stating

that Firth's estimator often performs best across circumstances. Rainey and McCaskey (2021) demonstrated that the penalized maximum likelihood estimator reduces both bias and variance when compared to standard MLE, yielding significant improvements in small samples (e.g., 50 observations) and perceptible gains even in larger datasets (e.g., 1,000 observations). Zorn (2005) supported the use of Firth's logistic regression in political science, stressing its advantages for stratified datasets prone to separation. While recognizing its effectiveness, he observed that its application remained limited beyond biostatistics, indicating a need for broader disciplinary dissemination.

Collectively, the scholarly contributions establish a strong foundation for the current study. They illustrate the development and application of Firth penalized regression, its empirical advantages, and its disciplinary adoption. However, several gaps persist in the literature (De Oliveira et al., 2019). First, while the methodological effectiveness of Firth's correction is well-documented, the bibliographic and thematic evolution of its usage remains uncharted. Second, little is known about the global research network, citation patterns, and institutional contributions to this domain. Lastly, despite the increasing volume of studies employing Firth regression, no bibliometric study has yet attempted to comprehensively map the literature. The present study addresses these gaps by conducting a detailed bibliometric analysis of scholarly works on complete separation in logistic regression using Firth's penalized method. This study uses bibliometric tools and approaches to find publication trends, influential authors, leading journals, collaboration patterns, and emerging themes. This provides a meta-perspective on how the area has evolved and where future research prospects exist. The literature covered here not only informs the conceptual and methodological foundations of this work, but it also demonstrates its importance in furthering understanding of the scholarly effect and trend of Firth penalized regression. Fig. 1 depicts the phases of bibliometric analysis.

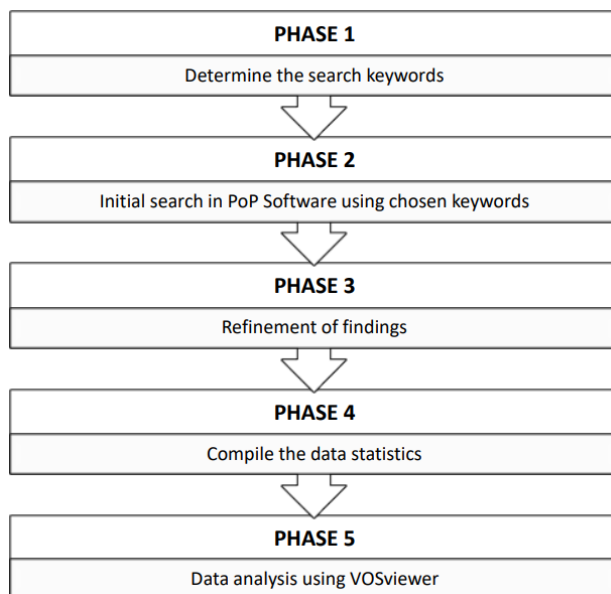


Fig. 1. The phase of bibliometric analysis of shift scheduling

Source: Sarudin et al. (2023)

3. METHODOLOGY

This study retrieved bibliographic data from the Scopus database on February 22, 2025. Scopus was selected as the primary database because it is one of the largest and most comprehensive abstract and

citation databases of peer-reviewed literature. It is widely recognized for its extensive multidisciplinary coverage and is frequently used in bibliometric and literature review studies (Rojas-Flores et al., 2023; Sarudin et al., 2024). Compared to other digital databases, Scopus has been found to offer the most diverse field coverage (Yaman et al., 2019). To ensure data integrity, records were cross-checked for duplicate entries by comparing titles, authors, and publication years.

Bibliographic data can also be retrieved using software such as Publish or Perish, developed by Anne-Will Harzing (Harzing, 2023). However, it imposes a limit of 1,000 records per query (Sarudin et al., 2023). This constraint reduces its suitability for comprehensive bibliometric reviews. Scopus, in contrast, allows for larger and more flexible data retrieval (Baas et al., 2020). Although Dimensions has emerged as a potential alternative to Scopus and Web of Science and often provides similar coverage (Harzing, 2019; Martín-Martín et al., 2021), it still showed reduced yield for this specific topic. Google Scholar was also excluded due to reproducibility limitations. In contrast, Scopus offers more stable and reproducible results across regions (Pozsgai et al., 2020), which is critical for ensuring the transparency and reliability of bibliometric analyses.

The data are obtained using the topic search query, as illustrated in Fig. 2. The initial search string used was TITLE-ABS-KEY ((“complete separation” OR “separation issue”) AND (“logistic regression”) AND (“Firth” OR “Firth's correction” OR “penalized regression”)). The results then underwent a filtering process based on several criteria including document type, year of publication, and language. The refined Scopus search string was TITLE-ABS-KEY ((“complete separation” OR “separation issue”) AND (“logistic regression”) AND (“Firth” OR “Firth's correction” OR “penalized regression”)) AND (LIMIT-TO (PUBYEAR , 2012) OR LIMIT-TO (PUBYEAR , 2017) OR LIMIT-TO (PUBYEAR , 2018) OR LIMIT-TO (PUBYEAR , 2019) OR LIMIT-TO (PUBYEAR , 2022) OR LIMIT-TO (PUBYEAR , 2023)) AND (LIMIT-TO (DOCTYPE , “ar”)) AND (LIMIT-TO (LANGUAGE , “English”))).

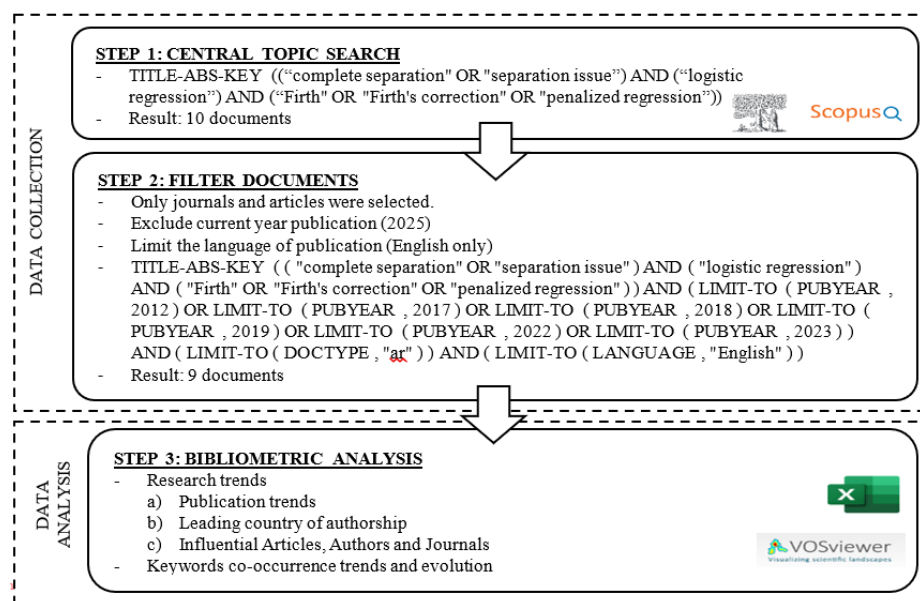


Fig. 2. Workflow of data retrieval and analysis for the bibliometric review

The inclusion criteria consisted of journal articles published in English between 2012 and 2024. Only English-language publications were included in the analysis to ensure consistency in text-mining, keyword analysis, and citation matching, as non-English articles often lack standardized metadata

in bibliographic databases. In addition, multiple investigations suggest that restricting reviews to English-language publications has minimal effect on overall conclusions (Nussbaumer-Streit et al., 2020; Dobrescu et al., 2021).

The remaining documents are then analyzed to determine research trends in leading countries. The analysis also identifies prominent authors and journals that actively publish articles related to complete separation in logistic regression using Firth penalized regression. Additionally, frequently used keywords from previous research were examined. Initially, ten documents on this topic were identified in the Scopus database, covering the period from 2012 to 2025. Of these, nine articles met the inclusion criteria of being journal articles published in English up to the year 2024 as suggested by Sarudin et al. (2023) and Ilmasari et al. (2022). These articles were then exported in Comma-Separated Values (CSV) format for bibliometric analysis.

This study employs several analytical tools to analyze and visualize the findings. VOSviewer will be used to conduct country co-authorship analysis and keyword co-occurrence analyses. It will also illustrate the results through network visualizations that show the connections between documents, authors, or countries in the dataset (Van Eck & Waltman, 2023; Ilmasari et al., 2022). The stronger the correlation between two nodes, the higher the link strength assigned (Ilmasari et al., 2022). According to Perianes-Rodriguez et al. (2016), VOSviewer offers two counting methods: full counting and fractional counting.

In VOSviewer, the country co-authorship analysis was conducted using full counting, where each co-author's country receives full credit for a publication regardless of the number of co-authors from other countries. Full counting is straightforward to interpret and provides an intuitive measure of total participation, making it particularly suitable for identifying absolute productivity and visibility of countries in a small dataset (Iskandar et al., 2020). An alternative is fractional counting, where each publication's credit is divided proportionally among co-authors' countries. It offers a more precise representation of relative contribution. It can underrepresent the role of countries engaged in large and highly collaborative projects (Donner, 2020).

In this study, the dataset contains nine publications and collaboration patterns are relatively sparse, full counting was preferred as it avoids diminishing the visibility of countries involved in multi-country papers. Besides, it assigns integer values to link strength based on the number of co-authored documents (Van Eck & Waltman, 2023). In addition to VOSviewer, Microsoft Excel will be used to compute data metrics and generate table.

4. FINDINGS

This bibliometric analysis presents results in two subsections: (1) research trends and (2) keyword evolution across chronological groupings. The findings aim to help future researchers understand the development of this area and identify opportunities for new topics or algorithms.

4.1 Research trends

The research trends related to complete separation in logistic regression using Firth penalized regression are examined in this section. The discussion covers publication trends, authorship by country, and influential articles, authors, and journals.

4.1.1 Publication Trends

Based on data obtained from the Scopus database, the number of publications on the topic of complete separation in logistic regression using Firth penalized regression has shown significant growth particularly over the past decade. This trend reflects a growing interest and ongoing advancements in the field. A total of ten documents published between 2012 and 2025 were identified where the topic appeared in the title,

abstract, or keywords. Among these, journal articles were the most common publication type, accounting for all ten documents. After the screening phase, only English-language journal articles published up to the year 2024 were selected, resulting in a final dataset of nine journal articles. All nine articles included in this study are published in peer-reviewed, properly archived journals indexed in Scopus, ensuring scholarly quality and long-term accessibility. This improves the reliability of the sources and the strength of the bibliometric mapping.

As shown in Fig. 3, the highest number of publications occurred in 2023, 2022, and 2019, with two articles each. In contrast, only one article was published in 2018, 2017, and 2012. An analysis of the publication trends on complete separation in logistic regression using Firth penalized regression reveals a clear pattern across two distinct phases. In the initial phase (2012–2018), only three articles were published, accounting for 33.33% of the total. In the development phase (2019–2024), the number increased to six articles, representing 66.67% of the publications. This growth indicates an increasing research interest and a broader exploration of methodologies, indicating significant progress in the field.

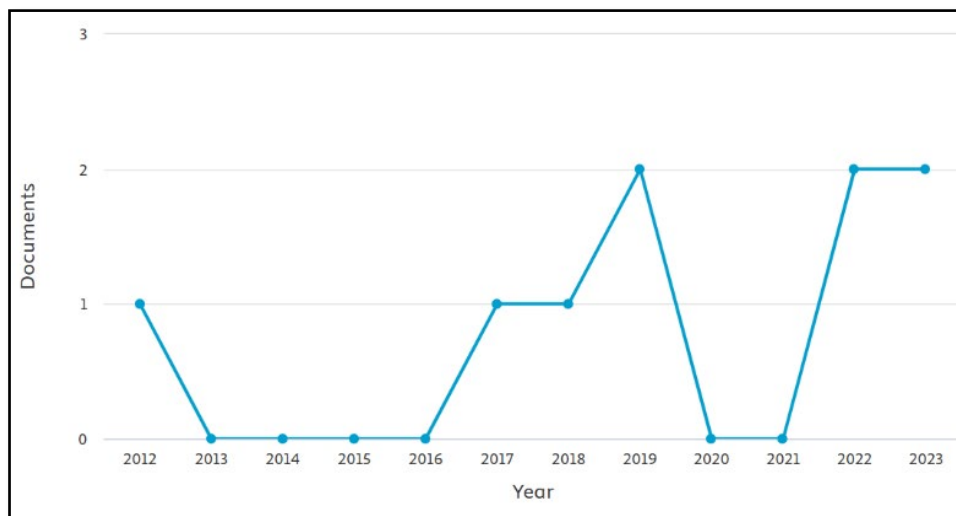


Fig. 3. Annual publication trends based on Scopus data, 2012–2024

4.1.2 Leading country of authorship

The geographical distribution of research on complete separation in logistic regression using Firth penalized regression demonstrates worldwide interest in this study area and highlights collaborative efforts. This subsection examines contributions from various countries, focusing on the leading countries with the most articles published in the Scopus database. VOSviewer was used to analyze and visualize international collaboration through network maps. The co-authorship analysis was conducted with full counting, considering only countries with at least one article and a maximum of 25 countries per document.

As a result, eight countries published articles on complete separation in logistic regression using Firth penalized regression between 2012 and 2024. The United States contributed two articles (Noor & Asmael, 2023; Choi et al., 2018), while Malaysia also produced two articles (Abdullah et al., 2022a; Abdullah et al., 2022b), making them the most productive countries in this area. Malaysia began its contributions in 2022, while the United States started in 2018. In addition, Slovenia, India, South Korea, Poland, Iraq, and Australia each contributed one article. Beyond focusing the leading countries, this study also explores

international collaboration in this research domain. International collaborations among these eight countries were analyzed using VOSviewer and are visualized in a network map (Fig. 4). All eight countries met the inclusion threshold at least one document per country and a maximum of 25 countries per document.

In the network, the country co-authorship network based on the identified publications. Each node represents a country, and the spatial distribution suggests limited collaboration among countries as no visible links (edges) are present between the nodes. This indicates that the publications in the dataset were largely authored within single countries without international co-authorship. Countries such as the United States and Malaysia appear more prominently, which may reflect a higher number of publications rather than actual collaboration. In contrast, countries like Iraq, India, Poland, and South Korea are represented as isolated nodes, suggesting opportunities for strengthening international research ties in this area.

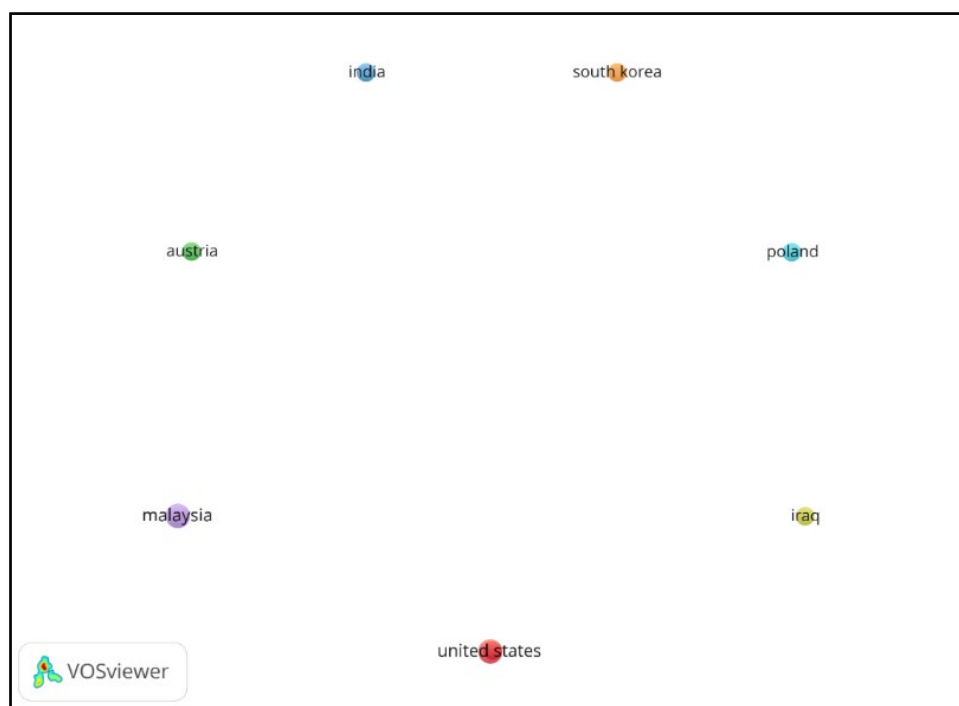


Fig. 4. Network visualization map of international collaboration on complete separation in logistic regression using Firth penalized regression research from 2012 to 2024.

Source: Online map: <https://tinyurl.com/27vfxvfs>

These countries were clustered into groups based on the frequency of collaboration using the Louvain algorithm for modularity optimization (Van Eck & Waltman, 2017; Blondel et al., 2008). This method detects communities by maximizing modularity, grouping together countries that collaborate more frequently with each other than with those outside the group. The number of clusters is not pre-defined. It is determined automatically by the algorithm based on the underlying network structure (Blondel et al., 2024; Rahiminejad et al., 2019).

In this dataset, the algorithm produced seven distinct clusters which reflects the relatively sparse collaboration patterns among countries in this research area. Countries that collaborate frequently tend to have similar co-authorship patterns (Sarudin et al., 2024). The clustering results were as follows: Cluster 1

(red) included Australia and the United States, indicating academic collaboration between them. Cluster 2 (green) consisted of Austria, Cluster 3 (blue) included only India, Cluster 4 (yellow) included Iraq, Cluster 5 (purple) consisted of Malaysia, Cluster 6 (sea blue) included Poland, and Cluster 7 (orange) consisted of South Korea.

4.1.3 Influential articles, authors, and journals

Table 1 presents the ranking of articles related to complete separation in logistic regression using Firth penalized regression based on the total number of citations. The highest-ranked article is "Maximum Likelihood and Firth Logistic Regression of the Pedestrian Route Choice," published in 2017 in the *International Regional Science Review* with a total of 39 citations, contributing approximately 38.61% to the overall citation count. This study authored by Gim and Ko (2016), suggests a significant impact in this research area. In the context of specialized statistical methodology studies where speciality topics often attract fewer total citations than broader applied research, accumulating nearly 40 citations over seven years can be considered relatively high.

The second highest-ranked article is "Bring more data! — A good advice? Removing separation in logistic regression by increasing sample size," published in 2019 in the *International Journal of Environmental Research and Public Health* with 18 citations and approximately 17.82% contribution. This article authored by Šinkovec et al. (2019), also demonstrates substantial influence on subsequent research. Other significant contributions include the article "Firth's penalized logistic regression: A superior approach for analysis of data from India's National Mental Health Survey, 2016," published in 2023 in the *Indian Journal of Psychiatry* with 13 citations and about 12.87% contribution of the total, authored by Suhas et al. (2023).

Another significant article is "Evaluating statistical approaches to leverage large clinical datasets for uncovering therapeutic and adverse medication effects," published in 2018 in *Bioinformatics* with 12 citations and approximately 11.88% contribution. It was authored by Choi et al. (2018). Interestingly, several recent articles have already gained considerable attention, suggesting the growing relevance of research on complete separation in logistic regression. Conversely, the article "A Study on Interstate Freight Mode Choice Between Trucks and Trains Used to Transport Oil Products: A Case Study of Iraq," published in 2023 in *Transport Problems* by Noor and Asmael (2023) has yet to accumulate citations at the time of analysis. This may be due to its recent publication date or narrower topical scope. It may require more time to gain recognition within the academic community.

Some journals within the dataset hold particularly strong influence in their respective fields. For example, *Bioinformatics* is a leading outlet in computational biology and bioinformatics, with a 2024 Impact Factor (IF) of 5.4 (Q1) and an SCImago Journal Rank (SJR) of approximately 2.45. These metrics reflect its high visibility and academic prominence. Similarly, the *Journal of Statistical Software* is widely regarded in the field of statistical methodology, with a 2024–2025 Impact Score of approximately 5.8, an SJR of 2.72, and a Q1 ranking, underscoring its methodological impact and widespread citation.

Other journals, such as the *International Journal of Environmental Research and Public Health*, *Australian Veterinary Journal*, and *International Regional Science Review*, also demonstrate moderate to strong standing in their respective domains, typically ranking in Q2 or Q3 quartiles. This mix of high- and mid-tier outlets indicates that the selected articles are drawn from both methodologically influential and applied research venues, providing balanced perspective on the scholarly landscape of research on complete separation in logistic regression using Firth penalized regression.

Table 1. The most influential article on complete separation in logistic regression using Firth penalized regression

No.	Title of article	Year	Journal	Total citations	Contribution rate (%)
-----	------------------	------	---------	-----------------	-----------------------

1.	Maximum Likelihood and Firth Logistic Regression of the Pedestrian Route Choice (Gim & Ko, 2016)	2017	International Regional Science Review	39	38.61
2.	Bring More Data! —A Good Advice? Removing Separation in Logistic Regression by Increasing Sample Size (Šinkovec et al., 2019)	2019	International Journal of Environmental Research and Public Health	18	17.82
3.	Firth's Penalized Logistic Regression: A Superior Approach for Analysis of Data from India's National Mental Health Survey, 2016 (Suhas et al., 2023)	2023	Indian Journal of Psychiatry	13	12.87
4.	Evaluating Statistical Approaches to Leverage Large Clinical Datasets for Uncovering Therapeutic and Adverse Medication Effects (Choi et al., 2018)	2018	Bioinformatics	12	11.88
5.	Separation-Resistant and Bias-Reduced Logistic Regression: STATISTICA Macro (Fijorek & Sokolowski, 2012)	2012	Journal of Statistical Software	10	9.90
6.	Identification of Blood-Based Multi-Omics Biomarkers for Alzheimer's Disease Using Firth's Logistic Regression (Abdullah et al., 2022a)	2022	Pertanika Journal of Science and Technology	6	5.94
7.	A Review of 91 Canine and Feline Red-Bellied Black Snake (Pseudechis Porphyriacus) Envenomation Cases and Lessons for Improved Management (Wun et al., 2022)	2022	Australian Veterinary Journal	2	1.98
8.	Discovering Potential Blood-Based Cytokine Biomarkers for Alzheimer's Disease Using Firth Logistic Regression (Abdullah et al., 2022b)	2019	Epidemiology Biostatistics and Public Health	1	0.99
9.	A Study on Interstate Freight Mode Choice Between Trucks and Trains Used to Transport Oil Products: A Case Study of Iraq (Noor & Asmael, 2023)	2023	Transport Problems	0	0.00

4.2 Keyword co-occurrence trends and evolution

The analysis explores the keyword co-occurrences in research on complete separation in logistic regression using VOSviewer. The analysis included all keywords identified by the author or indexed in the database. This clustering facilitates the identification of thematic groupings within the field, outlining the primary research areas and methodological focus present in the existing literature. Only documents containing at least two occurrences of a keyword were considered. As a result, 12 out of 151 keywords met the threshold and were visualized in Fig. 3. These twelve keywords were grouped into three distinct clusters as shown in Table 2.

In VOSviewer, total link strength (TLS) represents the overall strength of co-occurrence links between a given keyword (node) and all other keywords in the network. A higher TLS indicates that a keyword has stronger or more frequent associations with other keywords, reflecting its centrality and importance in the research field. TLS is calculated by summing the co-occurrence frequencies of the keyword with all others in the dataset (Van Eck & Waltman, 2017). In this study, keywords with higher TLS are considered more influential within the network as they are more strongly connected to multiple thematic areas.

The keyword “human” being the most frequently occurring keyword with three occurrences and a TLS of 13. This was followed by “logistic regression analysis” and “article,” both with three occurrences and a TLS of 11. Additionally, “logistic models,” “statistical model,” and “humans” each appeared twice, also with a TLS of 11. This highlights that research on complete separation in logistic regression most often focuses on applications involving human studies.

The prominence of the keyword “human” indicates that research on complete separation in logistic regression is most commonly applied in studies involving human subjects. This is consistent with its

relevance in biomedical, psychological, and clinical research where issues of separation frequently arise due to rare outcomes or sparse data structures (Mansournia et al., 2017). Such conditions make traditional logistic regression unreliable, thereby necessitating bias-reduction methods such as Firth's penalized likelihood. This observation is further supported by findings that Firth's correction is particularly effective in small or imbalanced human-subject datasets (Alam et al., 2022). These findings highlight the potential for broader adoption of this method in health-related applications.

Cluster 1 which exhibit strong co-occurrence links to both methodological and application-related terms, giving it high degree centrality in the network. It contains six keywords: "human," "logistic regression analysis," "article," "logistic models," "statistical model," and "humans". "Human" stands out as the most influential keyword with three occurrences and a high TLS of 13, showing strong connections with other terms in the cluster. Both "logistic regression analysis" and "article" follow closely with strong connections, each with three occurrences and a TLS of 11. The keywords "logistic models," "statistical model," and "humans" each occur twice and also have a TLS of 11, indicating their significant in the research network.

Cluster 2 includes four keywords: "biomarkers," "complete separation," "Firth logistic regression," and "logistic regression". Within this group, "Firth logistic regression" and "complete separation" are closely linked, highlighting their relevance in addressing separation issues in logistic regression models. Cluster 3 contains two keywords: "regression analysis" and "estimation method". This cluster emphasizes the methodological aspects of logistic regression, focusing on estimation techniques and analytical methods.

Table 2. Clusters of keywords and their total link strength (TLS)

Cluster	Keyword	Occurrences	Total Link Strength
Cluster 1 (Red)	Human	3	13
	Logistic regression analysis	3	11
	Article	3	11
	Logistic models	2	11
	Statistical model	2	11
	Humans	2	11
Cluster 2 (Green)	Logistic regression	2	9
	Firth logistic regression	3	4
	Complete separation	3	3
	Biomarkers	2	2
Cluster 3 (Blue)	Regression analysis	2	10
	Estimation method	2	10

Fig. 6 presents a network visualization of keyword co-occurrences related to complete separation in logistic regression using VOSviewer. The analysis includes 11 keywords that met the minimum threshold of two occurrences. The network map organizes the keywords into three distinct clusters based on their co-occurrence patterns and total link strength.

The thickness of the lines connecting the nodes represents the strength of the relationships between the keywords (Sarudin et al., 2024). For instance, "human" shows strong connections with other keywords within Cluster 1, indicating that research on logistic regression in human-related data is a major focus. In contrast, the relatively weaker links between Cluster 2 and Cluster 3 suggest that methodological developments such as the Firth correction are less connected to applications involving human related data. This indicates a potential research gap. Such gaps from this network map can be determined by examining the links between nodes. If two keywords lack a connection, it suggests a promising area for future exploration in complete separation in logistic regression research.

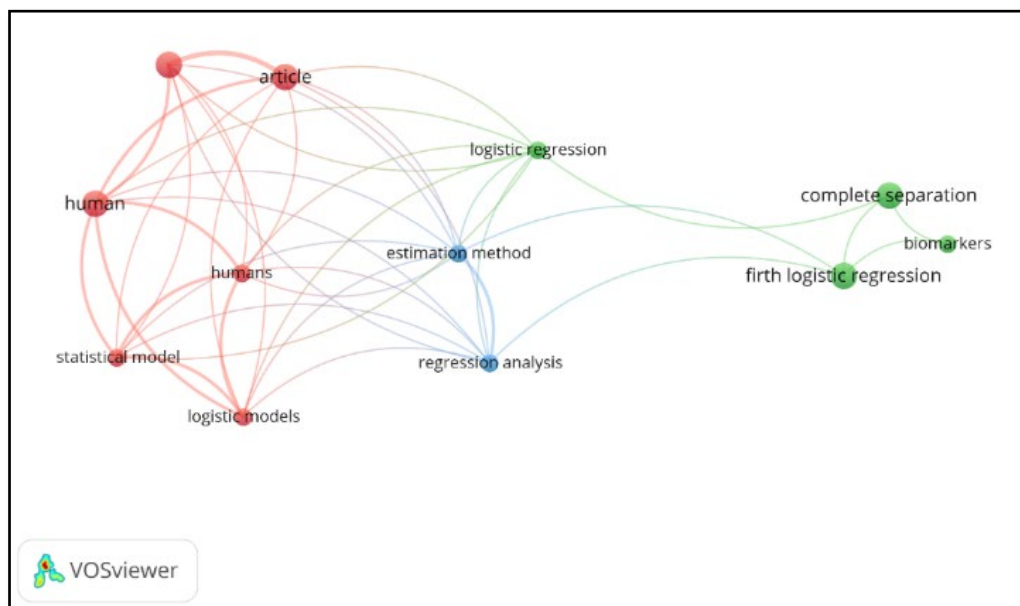


Fig. 6. Network node map of the exploration phase of complete separation in logistic regression research from 2012-2024.

Source: Online map: <https://tinyurl.com/2brypyl9>

5. CONCLUSION

This study presents the bibliometric analysis of research trends on complete separation in logistic regression using Firth penalized regression. By assessing nine selected articles published between 2012 and 2024, the study reveals a growing scholarly interest, particularly after 2019 with the United States and Malaysia emerging as the most productive contributors. Based on VOSviewer's co-authorship analysis, the United States also showed the highest collaboration connectivity, while Malaysia demonstrated strong regional collaboration despite its recent entry into the field. Influential articles primarily advanced methodological frameworks and applied solutions in health and transportation research. Keyword co-occurrence analysis identified main thematic clusters, such as human studies, statistical modelling, and estimation techniques, while also revealing underexplored areas. Although the small dataset reflects the niche and emerging nature of this research, the bibliometric analysis revealed significant trends in collaborations, publications, and thematic keyword evolution. The study acknowledges limitations such as the broad nature of some keywords such as 'human' and 'article', suggesting keyword normalization techniques for future research to improve thematic specificity. These findings lay a foundation for expanding interdisciplinary applications and addressing research gaps. Ultimately, this work not only enhances understanding of the academic landscape but also paves the way for methodological advancements and broader applications of Firth's approach in logistic regression.

6. ACKNOWLEDGEMENTS/FUNDING

The authors gratefully acknowledge the Faculty of Computer and Mathematical Sciences (FSKM) of Universiti Teknologi MARA (UiTM) at both the Tapah Campus and Shah Alam Campus for their institutional support and research facilities which made this work possible.

<https://doi.org/10.24191/jcrinn.v10i2.535>

7. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

8. AUTHORS' CONTRIBUTIONS

Nurul Husna Jamian: Study conception and design, Methodology development, Data analysis, Investigation, Original draft preparation. **Ahmad Zia Ul-Saufie Mohamad Japeri:** Supervision, Results validation, Manuscript review. **Mohammad Nasir Abdullah:** Research resources provision, Supervision, Manuscript review and editing.

9. REFERENCES

- Abdullah, M. N., Wah, Y. B., Majeed, A. B. A., Zakaria, Y., & Shaadan, N. (2022a). Identification of blood-based multi-omics biomarkers for Alzheimer's disease using Firth's logistic regression. *Pertanika Journal of Science & Technology*, 30(2), 1197–1218. <https://doi.org/10.47836/pjst.30.2.19>
- Abdullah, M. N., Wah, Y. B., Zakaria, Y., Majeed, A. B. A., & Huat, O. S. (2022b). Discovering potential blood-based cytokine biomarkers for Alzheimer's disease using Firth logistic regression. *Epidemiology Biostatistics and Public Health*, 16(4). <https://doi.org/10.2427/13173>
- Alam, T. F., Rahman, M. S., & Bari, W. (2022). On estimation for accelerated failure time models with small or rare event survival data. *BMC Medical Research Methodology*, 22, 169. <https://doi.org/10.1186/s12874-022-01638-1>
- Allison, P. D. (2008). Convergence failures in logistic regression. In *SAS Global Forum 2008* (Vol. 360, No. 1, p. 11). <http://www2.sas.com/proceedings/forum2008/360-2008.pdf>
- Baas, J., Schotten, M., Plume, A., Côté, G., & Karimi, R. (2020). Scopus as a curated, high-quality bibliometric data source for academic research in quantitative science studies. *Quantitative Science Studies*, 1(1), 377–386. https://doi.org/10.1162/qss_a_00019
- Blondel, V., Guillaume, J. L., & Lambiotte, R. (2024). Fast unfolding of communities in large networks: 15 years later. *Journal of Statistical Mechanics Theory and Experiment*, 2024, 10R001. <https://doi.org/10.1088/1742-5468/ad6139>
- Blondel, V. D., Guillaume, J. L., Lambiotte, R., & Lefebvre, E. (2008). Fast unfolding of communities in large networks. *Journal of Statistical Mechanics: Theory and Experiment*, 2008, P10008. <https://doi.org/10.1088/1742-5468/2008/10/P10008>
- Botes, M., & Fletcher, L. (2014, January). Comparing logistic regression methods for a sparse data set when complete separation is present. In *Annual Proceedings of the South African Statistical Association Conference* (Vol. 2014, No. con-1, pp. 1-8). South African Statistical Association (SASA).
- Clark, R. G., Blanchard, W., Hui, F. K. C., Tian, R., & Woods, H. (2023). Dealing with complete separation and quasi-complete separation in logistic regression for linguistic data. *Research Methods in Applied Linguistics*, 2(1), 100044. <https://doi.org/10.1016/j.rmal.2023.100044>
- Choi, L., Carroll, R. J., Beck, C., Mosley, J. D., Roden, D. M., Denny, J. C., & Van Driest, S. L. (2018). Evaluating statistical approaches to leverage large clinical datasets for uncovering therapeutic and adverse medication effects. *Bioinformatics*, 34(17), 2988–2996. <https://doi.org/10.1093/bioinformatics/bty306>

- D'Angelo, G., & Ran, D. (2024). Tutorial on Firth's logistic regression models for biomarkers in preclinical space. *Pharmaceutical Statistics*, 24(1), e2422. <https://doi.org/10.1002/pst.2422>
- De Oliveira, O. J., Da Silva, F. F., Juliani, F., Barbosa, L. C. F. M., & Nunes, T. V. (2019). Bibliometric method for mapping the state-of-the-art and identifying research gaps and trends in literature: An essential instrument to support the development of scientific projects. *IntechOpen*. <https://doi.org/10.5772/intechopen.85856>
- Dobrescu, A., Nussbaumer-Streit, B., Klerings, I., Wagner, G., Persad, E., Sommer, I., Herkner, H., & Gartlehner, G. (2021). Restricting evidence syntheses of interventions to English-language publications is a viable methodological shortcut for most medical topics: A systematic review. *Journal of Clinical Epidemiology*, 137, 209–217. <https://doi.org/10.1016/j.jclinepi.2021.04.012>
- Donner, P. (2020). A validation of coauthorship credit models with empirical data from the contributions of PhD candidates. *Quantitative Science Studies*, 1(2), 551–564. https://doi.org/10.1162/qss_a_00048
- Fijorek, K., & Sokolowski, A. (2012). Separation-resistant and bias-reduced logistic regression: STATISTICA macro. *Journal of Statistical Software, Code Snippets*, 47(2), 1–12. <https://doi.org/10.18637/jss.v047.c02>
- Firth, D. (1993). Bias reduction of maximum likelihood estimates. *Biometrika*, 80(1), 27–38. <https://doi.org/10.1093/biomet/80.1.27>
- Gilbert, N. (2022). *Logistic regression* (ebook). Taylor & Francis Group.
- Gim, T. H. T., & Ko, J. (2016). Maximum likelihood and Firth logistic regression of the pedestrian route choice. *International Regional Science Review*, 40(6), 616–637. <https://doi.org/10.1177/0160017615626214>
- Harzing, A. W. (2019). Two new kids on the block: How do Crossref and Dimensions compare with Google Scholar, Microsoft Academic, Scopus and the Web of Science? *Scientometrics*, 120, 341–349. <https://doi.org/10.1007/s11192-019-03114-y>
- Harzing, A. W. (2023). *Publish or Perish user's manual*. Harzing.com. <https://harzing.com/resources/publish-or-perish/manual>
- Heinze, G., & Schemper, M. (2002). A solution to the problem of separation in logistic regression. *Statistics in Medicine*, 21(16), 2409–2419.
- Hess, A. S., & Hess, J. R. (2019). Logistic regression. *Transfusion*, 59(7), 2197–2198. <https://doi.org/10.1111/trf.15406>
- Ilmasari, D., Sahabudin, E., Riyadi, F. A., Abdullah, N., & Yuzir, A. (2022). Future trends and patterns in leachate biological treatment research from a bibliometric perspective. *Journal of Environmental Management*, 318, 115594. <https://doi.org/10.1016/j.jenvman.2022.115594>
- Iskandar, A., Azis, F., Dewi, R. D. C., Rusli, R., & Ahmar, A. S. (2020). Co-authorship visualization of research on COVID-19 from Web of science data using bibliometric analysis. *Library Philosophy and Practice (e-journal)*, 4528, 1–9. <https://digitalcommons.unl.edu/libphilprac/4528>
- Karabon, P. (2020, March). Rare events or non-convergence with a binary outcome? The power of Firth regression in PROC LOGISTIC. In *SAS Global Forum 2020*. <https://www.sas.com/content/dam/SAS/support/en/sas-global-forum-proceedings/2020/4654-2020.pdf>
- Kumar, R. (2025). Bibliometric analysis: Comprehensive insights into tools, techniques, applications, and <https://doi.org/10.24191/jcrinn.v10i2.535>

- solutions for research excellence. *Spectrum of Engineering and Management Sciences*, 3(1), 45–62. <https://doi.org/10.31181/sems31202535k>
- Mansournia, M. A., Geroldinger, A., Greenland, S., & Heinze, G. (2017). Separation in logistic regression: Causes, consequences, and control. *American Journal of Epidemiology*, 187(4), 864–870. <https://doi.org/10.1093/aje/kwx299>
- Martín-Martín, A., Thelwall, M., Orduna-Malea, E., & López-Cózar, E. D. (2021). Google scholar, Microsoft academic, Scopus, Dimensions, Web of science, and open citations' coci: A multidisciplinary comparison of coverage via citations. *Scientometrics*, 126, 871–906. <https://doi.org/10.1007/s11192-020-03690-4>
- Nathanson, B. H., & Higgins, T. L. (2008). An introduction to statistical methods used in binary outcome modeling. *Seminars in Cardiothoracic and Vascular Anesthesia*, 12(3), 153–166. <https://doi.org/10.1177/1089253208323415>
- Noor, H. M., & Asmael, N. M. (2023). A study on interstate freight mode choice between trucks and trains used to transport oil products: A case study of Iraq. *Transport Problems*, 18(3), 29–40. <https://doi.org/10.20858/tp.2023.18.3.03>
- Nussbaumer-Streit, B., Klerings, I., Dobrescu, A. I., Persad, E., Stevens, A., Garritty, C., Kamel, C., Affengruber, L., King, V. J., & Gartlehner, G. (2020). Excluding non-English publications from evidence-syntheses did not change conclusions: A meta-epidemiological study. *Journal of Clinical Epidemiology*, 118, 42–54. <https://doi.org/10.1016/j.jclinepi.2019.10.011>
- Olowe, K. J., Edoh, N. L., Zouo, S. J. C., & Olamijuwon, J. (2024). Comprehensive review of logistic regression techniques in predicting health outcomes and trends. *World Journal of Advanced Pharmaceutical and Life Sciences*, 7(2), 016–026. <https://doi.org/10.53346/wjapls.2024.7.2.0039>
- Perianes-Rodriguez, A., Waltman, L., & Van Eck, N. J. (2016). Constructing bibliometric networks: A comparison between full and fractional counting. *Journal of Informetrics*, 10(4), 1178–1195. <https://doi.org/10.1016/j.joi.2016.10.006>
- Pozsgai, G., Lövei, G. L., Vasseur, L., Gurr, G., Batáry, P., Korponai, J., Littlewood, N. A., Liu, J., Móra, A., Obrycki, J., Reynolds, O., Stockan, J. A., VanVolkenburg, H., Zhang, J., Zhou, W., & You, M. (2020). A comparative analysis reveals irreproducibility in searches of scientific literature. *bioRxiv (Cold Spring Harbor Laboratory)*. <https://doi.org/10.1101/2020.03.20.997783>
- Puhr, R., Heinze, G., Nold, M., Lusa, L., & Geroldinger, A. (2017). Firth's logistic regression with rare events: Accurate effect estimates and predictions? *Statistics In Medicine*, 36(14), 2302–2317. <https://doi.org/10.1002/sim.7273>
- Rahiminejad, S., Maurya, M. R., & Subramaniam, S. (2019). Topological and functional comparison of community detection algorithms in biological networks. *BMC Bioinformatics*, 20, 212. <https://doi.org/10.1186/s12859-019-2746-0>
- Rainey, C., & McCaskey, K. (2021). Estimating logit models with small samples. *Political Science Research and Methods*, 9(3), 549–564. <https://doi.org/10.1017/psrm.2021.9>
- Rojas-Flores, S., Ramirez-Asis, E., Delgado-Caramutti, J., Nazario-Naveda, R., Gallozzo-Cardenas, M., Diaz, F., & Delfin-Narcizo, D. (2023). An analysis of global trends from 1990 to 2022 of microbial fuel cells: A bibliometric analysis. *Sustainability*, 15(4), 3651. <https://doi.org/10.3390/su15043651>
- Sarudin, E. S., Ariffin, W. N. M., & Jamian, S. S. (2024). Mapping the landscape: A bibliometric analysis of staff scheduling optimization research trends and keywords evolution. *International Journal of Research and Innovation in Social Science*, 8(8), 358–372. <https://doi.org/10.47772/ijriss.2024.808029>

- Sarudin, E. S., Aziz, W. N. H. W. A., Saleh, S. S. M., & Arsad, R. (2023). An overview of bibliometric indices and keyword classification in shift scheduling. *International Journal of Academic Research in Economics and Management Sciences*, 12(2), 289-302. <https://doi.org/10.6007/ijarems/v12-i2/17317>
- Šinkovec, H., Geroldinger, A., & Heinze, G. (2019). Bring more data! —A good advice? Removing separation in logistic regression by increasing sample size. *International Journal of Environmental Research and Public Health*, 16(23), 4658. <https://doi.org/10.3390/ijerph16234658>
- Stolte, M., Herbrandt, S., & Ligges, U. (2024). A comprehensive review of bias reduction methods for logistic regression. *Statistics Surveys*, 18, 139-162. <https://doi.org/10.1214/24-ss148>
- Suhas, S., Manjunatha, N., Kumar, C. N., Benegal, V., Rao, G. N., Varghese, M., & Gururaj, G. (2023). Firth's penalized logistic regression: A superior approach for analysis of data from India's national mental health survey, 2016. *Indian Journal of Psychiatry*, 65(12), 1208–1213. https://doi.org/10.4103/indianjpsychiatry.indianjpsychiatry_827_23
- Van Eck, N. J., & Waltman, L. (2017). Citation-based clustering of publications using CitNetExplorer and VOSviewer. *Scientometrics*, 111, 1053–1070. <https://doi.org/10.1007/s11192-017-2300-7>
- Van Eck, N. J., & Waltman, L. (2023). *VOSviewer manual* (Version 1.6.19). Centre for Science and Technology Studies, Leiden University. https://www.vosviewer.com/documentation/Manual_VOSviewer_1.6.19.pdf
- Walker, D. A., & Smith, T. J. (2019). Logistic regression under sparse data conditions. *Journal of Modern Applied Statistical Methods*, 18(2), eP3372. <https://doi.org/10.22237/jmasm/1604190660>
- Wun, M. K., Padula, A. M., Greer, R. M., & Leister, E. M. (2022). A review of 91 canine and feline red-bellied black snake (*pseudechis porphyriacus*) envenomation cases and lessons for improved management. *Australian Veterinary Journal*, 100(7), 318–328. <https://doi.org/10.1111/avj.13159>
- Yaman, A., Yoganingrum, A., Yaniasih, Y., & Riyanto, S. (2019). Tinjauan pustaka sistematis pada basis data pustaka digital: Tren riset, metodologi, dan coverage fields. *Jurnal Dokumentasi Dan Informasi*, 40(1), 1-20. <https://doi.org/10.14203/j.baca.v40i1.481>
- Zorn, C. (2005). A solution to separation in binary response models. *Political Analysis*, 13(2), 157-170. <https://doi.org/10.1093/pan/mpi009>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Navigating AI in Higher Education: Examining Over-Reliance and Plagiarism among UiTM Tapah Students

Nor Aslily Sarkam^{1*}, Nor Hazlina Mohammad², Nurizatie Farhanim Saiful Lizam³, Nurul Ain Azmi⁴, Nadhira Yasmin Mazli⁵, Mizan Qamalia Mohd Dzahir⁶, Ros Amira Arisha Md Isa⁷

^{1,2,3,4,5,6,7}Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Perak Branch, Tapah Campus, 35400 Tapah Road, Perak, Malaysia.

ARTICLE INFO

Article history:

Received 9 May 2025

Revised 6 August 2025

Accepted 15 August 2025

Online first

Published 1 September 2025

Keywords:

Artificial Intelligence (AI)

Higher Education

Digital Literacy

Academic Integrity

Technology Adoption

DOI:

10.24191/jcrinn.v10i2

ABSTRACT

The rapid advancement of Artificial Intelligence (AI) has transformed various sectors, including higher education. This study explores the challenges and opportunities of AI tool adoption in higher education, focusing on student learning experiences and ethical considerations. By employing correlation and regression analysis, this research analyzes data from 357 students to examine key factors influencing AI adoption, including effort expectancy, performance expectancy, digital literacy, and behavioural intention. The findings suggest that digital literacy significantly affects students' acceptance of AI tools, reinforcing the importance of targeted educational interventions. While AI integration enhances learning efficiency and accessibility, ethical concerns such as data privacy, algorithmic bias, and academic integrity remain critical challenges. The study provides insights for educators, policymakers, and institutions to develop strategies that balance technological advancements with ethical responsibility, ensuring an inclusive and effective AI-driven educational environment. Future research should explore longitudinal impacts and cross-cultural variations in AI adoption.

1. INTRODUCTION

Artificial intelligence (AI) has become an essential component of modern society, influencing fields such as medical, business, and education. AI-powered systems can execute activities that need human intellect, such as disease diagnosis, music composition, hyper-realistic image generation, and even customer service using chatbots (Kushmar et al., 2022). In education, AI has the potential to improve learning experiences by fostering student autonomy, increasing creativity, and reducing instructor workload. Furthermore, AI can provide personalised learning possibilities, allowing students to investigate topics relevant to their abilities and needs (Gocen & Aydemir, 2020). UNESCO emphasises that AI literacy is critical for the twenty-first century since it not only improves global economic competitiveness but also prepares students for future job markets (Yim, 2024).

^{1*} Corresponding author. E-mail address: norasilysarkam@uitm.edu.com
<https://doi.org/10.24191/jcrinn.v10i2.529>

The growing dependence on AI in education, while beneficial, raises concerns about its effects on students' cognitive development and academic integrity. AI-generated recommendations can exhibit bias, mislead users, or present inaccuracies, potentially resulting in unethical decision-making or academic dishonesty (Spatola, 2024). The accessibility of AI-generated content has raised concerns regarding plagiarism, as students may utilise such text without appropriate attribution, thereby compromising academic integrity. Large Language Models (LLMs) such as ChatGPT produce text that closely resembles human writing, complicating the differentiation between original content and that generated with AI assistance. Consequently, institutions face challenges in evaluating the authenticity of student assignments, and AI-detection tools like Turnitin frequently do not accurately identify content generated by AI (Hutson, 2024).

A report from Common Sense Media reveals that 70% of teenagers have used AI tools to assist with schoolwork, with 40% doing so without their teacher's knowledge. Additionally, 37% of students report that their schools lack clear AI-related guidelines, while 42% mention that their teachers prohibit AI use (Fisher, 2024). Malaysian Prime Minister Datuk Seri Anwar Ibrahim has acknowledged the importance of AI in education but warned that excessive dependence on AI-generated responses could hinder students' critical thinking skills (Latif, 2024). Furthermore, AI has contributed to a rise in academic misconduct, as students can easily manipulate AI-generated content, giving unfair advantages to those who do not adhere to academic honesty (Ahmad & Rahman, 2024).

Given these concerns, this study aims to investigate the impact of AI on education, focusing on two major challenges: over-reliance on AI tools and plagiarism. The objectives of this study are:

- (i) To identify the factors influencing the impact of AI on education.
- (ii) To determine the relationship between plagiarism and over-reliance on the impact of AI on education.

By examining these issues, this study aims to provide insights into the appropriate integration of AI in education, ensuring that students may effectively use AI tools while keeping academic integrity and independent critical thinking abilities.

2. LITERATURE REVIEW

AI has rapidly conquered almost every aspect of human life, starting from education to exploration into space. While AI offers a great potential revolution in teaching and learning, integration of AI at the higher education level raises concerns around academic integrity. Plagiarism and over-reliance on AI tools have emerged as major challenges that seriously threaten the authenticity of students' work. This section discusses some of the earlier research that was conducted on plagiarism and over-reliance, and the influence of AI on higher education

2.1 Plagiarism towards impact of AI on education

The rise of generative artificial intelligence (AI), particularly large language models (LLMs) like ChatGPT, has significantly transformed academic writing and raised concerns about plagiarism and intellectual property (Hutson, 2024). As AI becomes increasingly embedded in academia, it necessitates a re-evaluation of originality, research ethics, and the frameworks governing intellectual property and plagiarism. A report by Turnitin's Chief Product Officer, Chechitelli (2023) found that out of 38.5 million student submissions analyzed for AI-generated content, 9.6% contained over 20% AI-generated text, while 3.5% exhibited between 80% and 100% AI content. These findings highlight the growing concern over AI-assisted plagiarism, often termed "aigiarism."

The unethical use of AI tools in academic writing has sparked debates regarding academic dishonesty. Khalaf (2025) identified a significant positive correlation between plagiarism and aigiarism, with linear regression analysis indicating that students' attitudes toward traditional plagiarism significantly predicted their acceptance of AI-assisted plagiarism. Furthermore, frequency analysis revealed that 27% of students had positive attitudes toward traditional plagiarism, while 57% viewed aigiarism favorably. These results suggest that while AI tools provide convenience, they also pose ethical dilemmas that need to be addressed.

Concerns about originality and academic integrity influence students' decisions regarding AI use. Malik et al. (2023) found that 86% of students refrained from using AI in essay writing due to fears of losing originality and critical thinking skills. Other concerns included misinformation and inaccuracies (70%) and the ethical implications of unintentional plagiarism (69%). These findings indicate that while AI can serve as a powerful aid in learning, students remain cautious about its potential drawbacks.

The rapid adoption of AI tools in academic settings has also raised alarms among educators. Gruenhagen et al. (2024) reported that over one-third of students have used AI chatbots for assessments, often without perceiving it as a violation of academic integrity. Similarly, Niloy et al. (2024) found a significant relationship between students' motivations for using AI chatbots and a decline in academic honesty, prompting ongoing debates within the scientific community regarding the ethical implications of AI-assisted writing.

Despite these concerns, AI remains a valuable tool for education and research. Livberber and Ayvaz (2023) highlight AI's ability to generate novel research ideas, simplify complex concepts, and improve the quality of academic work. However, ethical challenges such as plagiarism and misinformation must be carefully managed. Sozon et al. (2024) identified multiple factors contributing to plagiarism, including academic pressure, tight deadlines, peer influence, and low awareness of academic integrity.

Academic pressure significantly influences plagiarism tendencies. Atmini et al. (2024) found that students experiencing high academic stress were more likely to engage in plagiarism, driven by competition, language barriers, and fear of low grades. Luo and Kong (2024) examined AI's role in rewriting academic content and found that AI-assisted rewrites sometimes increased duplication rates rather than reducing them, raising concerns about the effectiveness of AI in maintaining originality.

In conclusion, while AI tools offer valuable support in academic writing, their misuse can lead to ethical concerns such as plagiarism and academic dishonesty. Educators and institutions must establish clear guidelines and emphasize ethical AI practices to ensure that AI enhances learning rather than compromises academic integrity.

2.2 Over-reliance toward impact of AI on education

The increasing integration of AI in education has raised concerns regarding students' over-reliance on AI-generated assistance and its implications for learning outcomes. Klingbeil et al. (2024) found that simply knowing advice is AI-generated increases users' trust in it, sometimes leading to poor decision-making and unintended consequences. Particularly in high-stakes situations, such as financial or ethical matters, individuals tend to rely on AI-generated recommendations even when they contradict other sources. The authors suggest designing AI systems that encourage balanced trust while educating users on AI's limitations to promote informed decision-making.

Similarly, Ahmad et al. (2023) examined AI's impact on students' decision-making, motivation, and privacy concerns in Pakistan and China. Their study of 285 university students revealed that increased AI usage was associated with a decline in decision-making skills, greater dependency on automation, and heightened privacy risks. Although AI enhances learning, excessive reliance may erode essential skills like critical thinking, prompting the need for responsible AI engagement.

In Malaysia, Ismail (2024) investigated university students' use of AI tools for academic writing. Their study, based on surveys and interviews with 182 students, found that 42% preferred Google Translate, 32% used ChatGPT, 14% relied on Quillbot, and 10% used Grammarly for language correction. While these tools improve writing quality, the study warns that excessive dependence can compromise accuracy, privacy, and writing proficiency. The researchers emphasize that AI should be an aid rather than a replacement for learning, encouraging students and educators to adopt a mindful approach.

Zhai et al. (2024) further explored the cognitive impact of students' reliance on AI-powered dialogue systems. Their review highlighted that while AI streamlines task and enhances efficiency, over-reliance can weaken cognitive skills such as decision-making, critical thinking, and problem-solving. In educational settings, where these skills are fundamental to academic success, the authors advocate for teaching students how to critically evaluate AI-generated content to prevent misinformation and excessive dependency.

By surveying 597 college students and applying latent profile analysis (LPA), Stojanov et al. (2024) were able to determine five unique user profiles for ChatGPT's educational purposes. These profiles included students who seldom ever used AI and others who relied on ChatGPT for all of their schoolwork. The results highlight the different levels of AI dependence; some students use AI as a supplement, while others give it complete control over their work, which raises questions about academic honesty and the integrity of their learning.

Research by Gruenhagen et al. (2024) on AI chatbots in Australian universities found that 79% of students use AI for some aspect of their education, including homework and tests. This finding lends credence to these worries. Although AI tools improve information retrieval and comprehension, there are worries about academic integrity due to their ubiquitous use. To encourage safe and ethical use of AI, the researchers suggest having students help create usage guidelines.

Malik et al. (2023) investigated the influence of AI tools on academic writing among 245 undergraduate students across 25 Indonesian tertiary institutions. Their findings revealed that AI positively impacted writing skills, self-confidence, and academic integrity awareness. However, concerns arose regarding AI's potential to diminish creativity and critical thinking. The study suggests a balanced approach to AI adoption, integrating educator perspectives to refine AI's role in academic settings.

Finally, Darvishi et al. (2024) examined AI's influence on student agency in peer feedback through an eight-week experimental study involving 1,625 students across 10 courses. Results indicated that students often relied on AI-generated feedback rather than actively engaging with learning processes. When AI assistance was removed, self-regulated strategies partially compensated, but they remained less effective. The authors conclude that while AI improves task quality, excessive reliance may hinder self-regulated learning (SRL), a key component of lifelong learning.

In conclusion, while AI tools enhance efficiency and learning, excessive dependence on them poses risks to critical human skills such as decision-making, critical thinking, and self-regulation. The reviewed studies highlight that unchecked reliance on AI may lead to diminished originality, reduced motivation, and academic dishonesty. Therefore, fostering a balanced approach through AI literacy, ethical guidelines, and educator involvement is crucial to ensuring AI serves as an aid rather than a substitute for essential cognitive skills.

3. METHODOLOGY

3.1 Research design

This study employs a quantitative research approach with a cross-sectional design to examine the relationship between AI usage, over-reliance, and plagiarism among students in higher education. A cross-sectional design is appropriate for this study as data is collected at a single point in time from a sample of

students to analyze existing patterns and relationships. This approach allows for an efficient assessment of students' perceptions and behaviors related to AI in education.

3.2 Population and sample size

The target population for this study consists of all 4,610 students enrolled at UiTM Perak Branch, Tapah. To determine an appropriate sample size, Cochran's formula (1977) was used to ensure a statistically reliable sample. The sample size was calculated as follows:

$$n = \frac{Z^2 pq}{d^2} \div \left(1 + \frac{Z^2 pq}{d^2 N} - 1 \right) \quad (1)$$

where,

n = required sample size

N = population size (4610 students at UiTM Perak Branch, Tapah)

Z = Z - score corresponding to the desired confidence level (1.96 for 95% confidence interval)

p = Estimated proportion of the population with the characteristic of interest (used 0.5)

$q = 1 - p$ (the proportion of the population without the characteristics, $Q=0.5$)

d = 0.05 of the margin error

3.3 Sampling design

A probability sampling technique ensures a fair representation of students across different faculties. Specifically, stratified sampling is employed by dividing the population into three strata based on faculties and selecting a proportionate number of students from each group. The distribution is as follows:

Table 1. Sample distribution for each faculty

Faculties	Frequency	Sample
Faculty of Applied Science	1405	108
Faculty of Accountancy	1639	126
Faculty of Computer and Mathematical Sciences	1566	120
	4610	357

3.4 Data collection method

Theoretical Framework Data was collected using an online questionnaire and distributed via Google Forms. The survey link was shared through digital communication platforms such as WhatsApp and Telegram, specifically targeting students from the three faculties at UiTM Perak Branch, Tapah. The questionnaire was structured into six sections. The first section included an introduction to the study along with an informed consent. The second section gathered demographic information. The third section focused on the impacts of AI use in education and contained four questions. The fourth section addressed students' over-reliance on AI, comprising seven questions. The fifth section consisted of five questions that assessed students' understanding and concerns about plagiarism. The final section included an appreciation note for their participation.

Section three to five were validated using a 7-point Likert scale, ranging from 1 (Strongly Disagree) to 7 (Strongly Agree). The questionnaire items were adapted from previous studies on technology acceptance and academic integrity by Davis (1989), Ajzen (1991), and Teo (2010), and were reviewed by experts for

content validity. The data collection process was conducted over a 14-day period from November 26, 2024, to December 9, 2024. The demographics information of the respondents is presented in Table 2.

Table 2. The students' demographics information

	Frequency (N)	Percentage (%)
Gender		
Male	171	47.9
Female	186	52.1
Age		
17-18	83	23.2
19-20	210	58.8
21-22	27	7.6
23-24	11	3.1
25 and above	26	7.3
Faculty		
Faculty of Applied Science	109	30.5
Faculty of Accountancy	128	35.9
College of Computing, Informatic & Mathematics	120	33.6
Semester		
1	30	8.4
2	20	5.6
3	79	22.1
4	53	14.8
5	150	42.0
6	7	2.0
7	18	5.0

3.5 Theoretical framework

Based on Fig. 1, it shows that the dependent variable was the impact of AI while the independent variables were plagiarism and over-reliance.

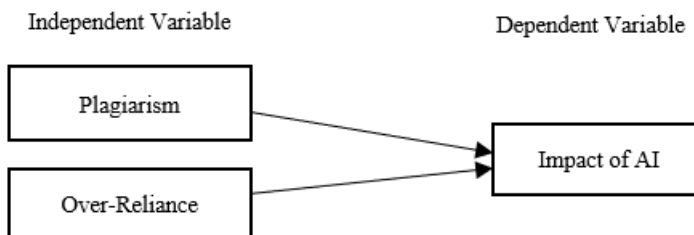


Fig. 1. Theoretical framework

3.6 Data preparation and analysis

Once data collection was completed, responses were compiled and transferred to Microsoft Excel for initial cleaning, where incomplete or inconsistent responses were checked and removed if necessary. The cleaned dataset was then imported into IBM SPSS Statistics 29 for further analysis.

The study employs descriptive analysis to summarize key demographics. Additionally, correlation and regression analyses were conducted to examine the relationship between students' over-reliance on AI and their engagement in plagiarism. Correlation analysis helps determine the strength and direction of

relationships between variables, while regression analysis provides insight into the predictive effects of AI dependence on plagiarism tendencies.

3.6.1 Correlation and regression analysis

To examine the relationship between over-reliance on AI and plagiarism among students, this study employs correlation and regression analysis. Correlation analysis measures the strength and direction of the relationship between two variables, while regression analysis determines the extent to which one variable influence another (Field, 2018).

Before conducting these analyses, several key assumptions must be met, including linearity, normality of residuals, homoscedasticity, and absence of multicollinearity (Tabachnick & Fidell, 2019). Linearity assumes a straight-line relationship between the independent and dependent variables. A violation of this assumption can lead to misleading conclusions, which can be assessed using scatter plots or Pearson correlation coefficients (Pallant, 2020). Normality of residuals requires that the residuals in the regression model follow a normal distribution. This assumption is critical for hypothesis testing and can be evaluated using histograms, P-P plots, or statistical tests such as the Shapiro-Wilk test (Hair et al., 2018). Homoscedasticity ensures that residuals exhibit constant variance across all levels of the independent variable, which can be tested using residual plots or the Breusch-Pagan test. A violation of homoscedasticity may result in inefficient estimates (Kline, 2016). Absence of multicollinearity is necessary to avoid highly correlated independent variables, which can distort regression estimates. Variance Inflation Factor (VIF) values exceeding 10 indicate severe multicollinearity, and corrective measures such as variable removal or Principal Component Analysis (PCA) can be applied (Gujarati et al., 2019).

The Pearson correlation coefficient (r) was used to quantify the strength of the relationship, where values closer to +1 or -1 indicate strong positive or negative correlations, respectively. For regression analysis, a multiple linear regression model was applied to assess how over-reliance on AI predicts plagiarism tendencies. The regression equation is:

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \varepsilon \quad (2)$$

where Y represents plagiarism, X represents over-reliance on AI, β_0 is the intercept, β_1 is the regression coefficient, and ε is the error term. By analyzing these relationships, this study provides empirical evidence on whether excessive AI usage significantly contributes to academic dishonesty, informing policies on AI integration in education.

4. RESULTS AND DISCUSSION

4.1 Reliability analysis

The reliability of the study's constructs was assessed using Cronbach's alpha. According to Nunnally (1978), a Cronbach's alpha value of 0.70 or higher indicates acceptable internal consistency, while values between 0.60 and 0.70 are considered moderate but acceptable in exploratory research. The results show that Over-Reliance (0.786) and Plagiarism (0.733) demonstrate good reliability, while Impact of AI (0.624) falls within the acceptable range (Taber, 2018). These findings suggest that the measurement scales used in this study are reliable for further analysis.

Table 3. Reliability analysis

Variables	Cronbach's Alpha
Impact of AI	0.624
Plagiarism	0.733
Over-Reliance	0.786

4.2 Correlation

Regression Analysis A Pearson correlation analysis was conducted to examine the relationships between over-reliance, plagiarism, and the impact of AI in higher education. The results, as shown in Table 4, indicate significant positive correlations among all variables at the 0.01 significance level (2-tailed).

Table 4. The correlation result

	Impact of AI	Plagiarism	Over Reliance
Impact of AI	1.00	0.502**	0.599**
Plagiarism		1.00	0.608**
Over Reliance			1.00

The correlation between over-reliance on AI and its impact is $r = .599$, $p < .001$, suggesting a moderate positive association. This indicates that students who exhibit a greater dependence on AI tend to perceive a higher impact of AI on their education. Similarly, plagiarism is significantly correlated with the impact of AI ($r = .502$, $p < .001$), demonstrating that higher levels of AI-assisted academic dishonesty are associated with a stronger perceived impact of AI in learning environments.

Additionally, there is a moderate correlation between over-reliance on AI and plagiarism ($r = .608$, $p < .001$), highlighting that students who excessively depend on AI tools are more likely to engage in academic dishonesty. These findings align with previous studies suggesting that the accessibility of AI-driven tools may lead to ethical challenges in academic settings (Çela et al. 2024).

The results suggest that AI reliance and plagiarism are interrelated factors influencing students' academic behavior and the perceived role of AI in education. Future research should explore strategies to mitigate these negative consequences while leveraging AI's benefits in learning.

4.3 Assumptions of the regression analysis

Multicollinearity was assessed using tolerance and variance inflation factor (VIF) values. According to Hair et al. (2018), tolerance values below 0.10 and VIF values above 10 indicate severe multicollinearity. In this study, both Plagiarism and Over-Reliance have tolerance values of 0.630 and VIF values of 1.588, suggesting no significant multicollinearity issues.

4.3.1 Multicollinearity

Multicollinearity was assessed using tolerance and variance inflation factor (VIF) values. According to Hair et al. (2018), tolerance values below 0.10 and VIF values above 10 indicate severe multicollinearity. In this study, both Plagiarism and Over-Reliance have tolerance values of 0.630 and VIF values of 1.588, suggesting no significant multicollinearity issues.

Table 5. Multicollinearity

Variables	Collinearity Statistics	
	Tolerance	VIF
Plagiarism	0.630	1.588
Over-Reliance	0.630	1.588

4.3.2 Normality

The normality of the data was assessed using both graphical and numerical methods. The q-q plot for the impact of AI variable indicates that data points closely follow the diagonal reference line, suggesting approximate normality. Additionally, skewness and kurtosis values were examined to support this assumption. According to Hair et al. (2018), skewness values between -1 and 1 and kurtosis values within ± 1 indicate an approximately normal distribution. the results show that impact of AI (-0.644, -0.113), plagiarism (-0.061, -0.061), and over-reliance (-0.532, -0.135) all fall within acceptable ranges, confirming normality (Tabachnick & Fidell, 2019). Since no severe deviations were observed, the normality assumption is satisfied, making the dataset suitable for parametric analyses such as correlation and regression.

Table 6. The normality result

Variables	Skewness	Kurtosis	Interpretation
Impact of AI	-0.644	-0.113	Approximate normal
Plagiarism	-0.061	-0.061	Approximate normal
Over Reliance	-0.532	-0.135	Approximate normal

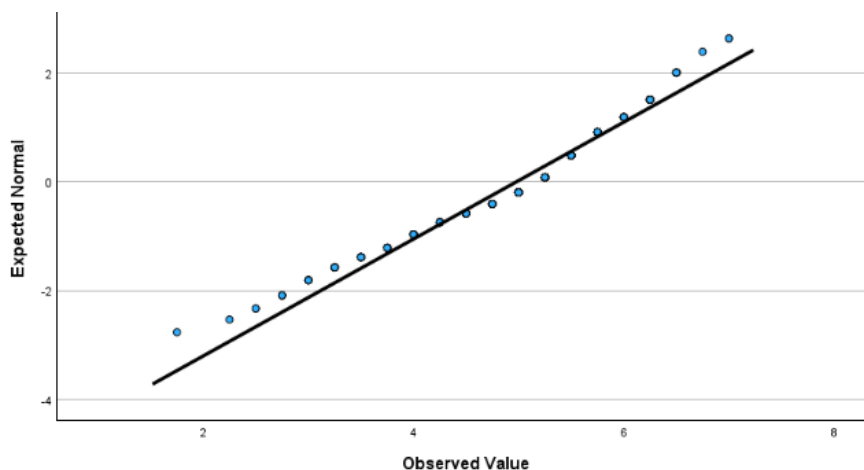


Fig. 2. Normal Q-Q plot

4.3.3 Homogeneity of variance

The scatterplot of standardized residuals versus standardized predicted values in Fig. 3 indicates no clear pattern, suggesting homoscedasticity, meaning the variance of errors remains constant across

predictions (Field, 2018). The residuals appear randomly dispersed, fulfilling the assumption of homogeneity of variance, which is essential for reliable regression analysis (Tabachnick & Fidell, 2019).

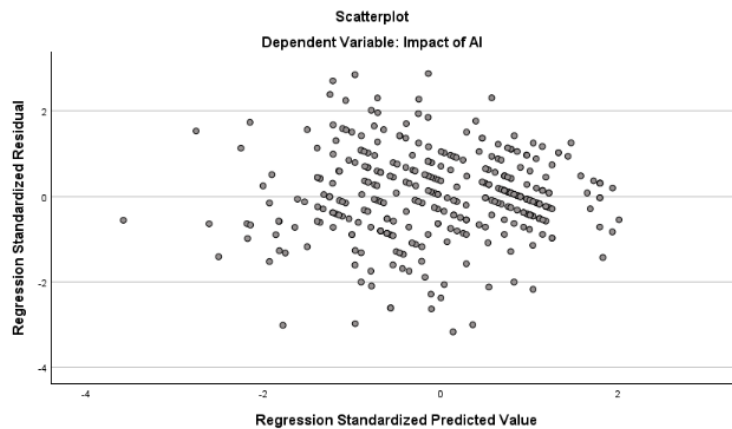


Fig. 3. Scatterplot of residual versus predicted

4.3.4 Linearity and normality of residuals

The assumptions of linearity were evaluated using a scatter plot of the observed versus predicted values. The plot revealed a reasonably straight-line pattern without notable curvilinear trends, indicating that the relationship between the independent variables (over-reliance on AI and plagiarism) and the dependent variable (the impact of AI) is approximately linear. This validates the appropriateness of applying linear regression to the data (Pallant, 2020).

Based on Fig. 2, the residuals are closely on the diagonal line, indicating that they are approximately normally distributed. This satisfies the normality assumption of regression, which is crucial for accurate p-values and confidence intervals in the model (Hair et al., 2018).

The Durbin Watson statistic, which was 1.958, further confirmed the independence of residuals, as values close to 2 suggest no autocorrelation (Field, 2018). Therefore, the assumptions of linearity, normality of residuals, and independence of errors were all sufficiently met.

4.4 Regression analysis

A multiple regression analysis was conducted to examine the influence of over-reliance on AI and plagiarism on the impact of AI in higher education. The model summary (Table 7) shows that the regression model is statistically significant, $R = .624$, $R^2 = .389$, Adjusted $R^2 = .386$, indicating that 38.9% of the variance in AI's impact is explained by the predictors ($F(2,354) = 112.740$, $p < .001$).

Table 7. Test of Significant of Coefficient

Variables	Unstandardized Coefficient Value	t-statistics	p-value
Constant	2.011	9.827	0.000
Plagiarism	0.435	8.925	0.000
Over - Reliance	0.210	4.152	0.000

The regression analysis reveals that both over-reliance on AI and plagiarism significantly predict the impact of AI in higher education. The regression equation is:

$$\text{Impact of AI} = 2.011 + 0.435 \text{ Plagiarism} + 0.210 \text{ Over} - \text{Reliance} \quad (3)$$

Both predictors show statistically significant relationships with the dependent variable, as indicated by their p-values (<0.001). Over-reliance on AI ($b=0.435$, $t = 8.925$, $p\text{-value} <0.001$) has a stronger influence compared to plagiarism ($b=0.210$, $t = 4.152$, $p\text{-value} <0.001$). These findings suggest that excessive dependence on AI and academic dishonesty contribute to its overall impact, supporting the need for policy interventions (Field, 2018).

5. CONCLUSION

The integration of AI tools in higher education presents both promising opportunities and notable challenges. This study highlights the significant role of digital literacy in shaping students' acceptance of AI-driven learning environments. Through the application of correlation and regression analysis, the research demonstrates that factors such as effort expectancy and performance expectancy directly impact students' behavioral intention to adopt AI tools. These findings align with existing literature, suggesting that technological competency and perceived usefulness are crucial determinants of AI acceptance (Venkatesh et al., 2003).

Despite the benefits of AI-enhanced learning, ethical concerns remain a major challenge. Issues such as data privacy, algorithmic bias, and the potential for academic dishonesty necessitate a balanced approach to AI implementation in educational settings. Institutions must establish clear guidelines to address these concerns, ensuring that AI tools are used responsibly and inclusively (Vieriu & Petrea, 2025).

Moreover, the study emphasizes the need for comprehensive digital literacy programs to bridge the gap between technological advancements and students' capabilities. Educators should be equipped with AI related training to facilitate effective implementation while mitigating ethical risks. Future research should explore the long-term effects of AI adoption on student learning outcomes and investigate cross-cultural perspectives on AI acceptance.

In conclusion, AI holds immense potential to revolutionize higher education, but its success depends on a well-structured framework that prioritizes ethical considerations and digital literacy. By fostering an inclusive AI-driven learning environment, higher education institutions can harness the full potential of AI while ensuring fairness, transparency, and accessibility for all students.

6. ACKNOWLEDGMENTS/FUNDING

The authors would like to acknowledge the support of Universiti Teknologi Mara (UiTM), Perak Branch, Tapah Campus for the support of the research support on this research.

7. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

8. AUTHORS' CONTRIBUTIONS

Nor Aslily Sarkam reviewed the interpretation of the results and supervised the overall writing of the article. **Nor Hazlina Mohammad** thoroughly proofreads the article and provided suggestions for improvement. **Nurizatie Farhanim Saiful Lizam**, **Nurul Ain Azmi**, and **Nadhira Yasmin Mazli** conducted the fieldwork, prepared the literature review and methodology. **Mizan Qamalia Mohd Dzahir** and **Ros Amira Arisha Md Isa** were responsible for instrument and data validation, reviewed the interpretation of the results, and contributed to refining the article.

9. REFERENCES

- Ahmad, N. A., & Rahman, D. (17 April, 2024). *AI dorong siswa ciplak tugasan, gugat etika akademik*. Berita Harian. <https://www.bharian.com.my/rencana/komentar/2024/04/1235730/ai-dorong-siswa-ci-plak-tugasan-gugat-etika-akademik>
- Ahmad, S. F., Han, H., Alam, M. M., Rehmat, M. K., Irshad, M., Arraño-Muñoz, M., & Ariza-Montes, A. (2023). Impact of artificial intelligence on human loss in decision making, laziness and safety in education. *Humanities and Social Sciences Communications*, 10, 311. <https://doi.org/10.1057/s41599-023-01787-8>
- Ajzen, I. (1991). The theory of planned behavior. *Organizational Behavior and Human Decision Processes*, 50(2), 179–211. [https://doi.org/10.1016/0749-5978\(91\)90020-T](https://doi.org/10.1016/0749-5978(91)90020-T)
- Atmini, S., Jusoh, R., Prastiwi, A., Wahyudi, S. T., Hardanti, K. N., & Widiarti, N. N. (2024). Plagiarism among accounting and business postgraduate students: A fraud diamond framework moderated by understanding of artificial intelligence. *Cogent Education*, 11(1), 2375077. <https://doi.org/10.1080/2331186X.2024.2375077>
- Çela, E., Fonkam, M. M., & Potluri, R. M. (2024). Risks of AI-assisted learning on student critical thinking: A case study of Albania. *International Journal of Risk and Contingency Management*, 12(1), 19. <https://doi.org/10.4018/IJRCM.350185>
- Chechitelli, A. (2023, May 30). *AI writing detection update from Turnitin's Chief Product Officer*. Turnitin. <https://www.turnitin.co.uk/blog/ai-writing-detection-update-from-turnitins-chief-product-officer>
- Darvishi, A., Khosravi, H., Sadiq, S., Gašević, D., & Siemens, G. (2024). Impact of AI assistance on student agency. *Computers and Education*, 210, 104967. <https://doi.org/10.1016/j.compedu.2023.104967>
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340. <https://doi.org/10.2307/249008>
- Field, A. (2018). *Discovering statistics using IBM SPSS statistics* (5th ed.). University of Sussex, UK, SAGE Publications Ltd.
- Fisher, J. (19 September 2024). *Parents, here's what you need to know about your teen's use of AI*. Lifewire. https://www.lifewire.com/teen-ai-use-8715109?utm_source=chatgpt.com
- Gocen, A., & Aydemir, F. (2020). Artificial intelligence in education and schools. *Research on Education and Media*, 12(1), 13–21. <https://doi.org/10.2478/rem-2020-0003>
- Gruenhagen, J. H., Sinclair, P. M., Carroll, J. A., Baker, P. R. A., Wilson, A., & Demant, D. (2024). The rapid rise of generative AI and its implications for academic integrity: Students' perceptions and use of chatbots for assistance with assessments. *Computers and Education: Artificial Intelligence*, 7,

100273. <https://doi.org/10.1016/j.cacai.2024.100273>
- Gujarati, D. N., Porter, D. C., & Pal, M. (2019). *Basic Econometrics* (6th ed.). MC Graw Hill India.
- Hair, J. F., Babin, B. J., Anderson, R. E., & Black, W. C. (2018). *Multivariate data analysis* (8th ed.). Cengage.
- Hutson, J. (2024). Rethinking plagiarism in the era of generative AI. *Journal of Intelligent Communication*, 3(2), 20-31. <https://doi.org/10.54963/jic.v3i2.220>
- Ismail, A. A. (2024). Over-reliance of AI tools in academic writing tasks: The Malaysian tertiary level context. In *E-Prosiding Persidangan Antarabangsa Sains Sosial & Kemanusiaan kali ke-9 (PASAK9 2024)* (pp. 218-224).
- Khalaf, M. A. (2025). Does attitude towards plagiarism predict aigiarism using ChatGPT? *AI and Ethics*, 5, 677-688. <https://doi.org/10.1007/s43681-024-00426-5>
- Kline, R. B. (2016). *Principles and practice of structural equation modeling* (4th ed.). The Guilford Press.
- Klingbeil, A., Grützner, C., & Schreck, P. (2024). Trust and reliance on AI — An experimental study on the extent and costs of overreliance on AI. *Computers in Human Behavior*, 160, 108352. <https://doi.org/10.1016/j.chb.2024.108352>
- Kushmar, L. V., Vornachev, A. O., Korobova, I. O., & Kaida, N. O. (2022). Artificial Intelligence in Language Learning: What Are We Afraid of. *Arab World English Journal (AWEJ) Special Issue on CALL(8)*. 262-273. <https://dx.doi.org/10.24093/awej/call8.18>
- Latif, F. (1 November 2024). *Rely too much on AI and we risk becoming powerless, PM Anwar warns students*. Malaymail. <https://www.malaymail.com/news/malaysia/2024/11/01/rely-too-much-on-ai-and-we-risk-becoming-powerless-pm-anwar-warns-students/155502>
- Livberber, T., & Ayvaz, S. (2023). The impact of Artificial Intelligence in academia: Views of Turkish academics on ChatGPT. *Heliyon*, 9(9), e19688. <https://doi.org/10.1016/j.heliyon.2023.e19688>
- Luo, H., & Kong, H. (2024). ChatGPT plagiarism in the academic field: Exploration and analysis of plagiarism effects. In *2024 International Conference on Machine Learning and Intelligent Computing*, (Vol. 245, pp. 1-11). Proceedings of Machine Learning Research.
- Malik, A. R., Pratiwi, Y., Andajani, K., Numertayasa, I. W., Suharti, S., Darwis, A., & Marzuki. (2023). Exploring artificial intelligence in academic essay: Higher education student's perspective. *International Journal of Educational Research Open*, 5, 100296. <https://doi.org/10.1016/j.ijedro.2023.100296>
- Niloy, A. C., Hafiz, R., Md, B., Hossain, B. M. T., Gulmeher, F., Sultana, N., Islam, K. F., Bushra, F., Islam, S., Hoque, S. I., Rahman, M. A., & Kabir, S. (2024). AI chatbots: A disguised enemy for academic integrity? *International Journal of Educational research Open*, 7, 100396. <https://doi.org/10.1016/j.ijedro.2024.100396>
- Nunnally, J. C. (1978). *Psychometric theory* (2nd ed.), McGraw Hill.
- Pallant, J. (2020). *SPSS survival manual: A step-by-step guide to data analysis using IBM SPSS* (7th ed.), Routledge. <https://doi.org/10.4324/9781003117452>
- Sozon, M., Pok, W. F., Sia, B. C., & Alkharabsheh, O. H. M. (2024). Cheating and plagiarism in higher education: a systematic literature review from a global perspective, 2016–2024. *Journal of Applied*

Research in Higher Education. <https://doi.org/10.1108/JARHE-12-2023-0558>

- Spatola, N. (2024). The efficiency-accountability tradeoff in AI integration: Effects on human performance and over-reliance. *Computers in Human Behavior: Artificial Humans*, 2(2), 100099. <https://doi.org/10.1016/j.chbah.2024.100099>
- Stojanov, A., Liu, Q., & Koh, J. H. L. (2024). University students' self-reported reliance on ChatGPT for learning: A latent profile analysis. *Computers and Education: Artificial Intelligence*, 6, 100243. <https://doi.org/10.1016/j.caeai.2024.100243>
- Tabachnick, B. G., & Fidell, L. S. (2019). *Using multivariate statistics* (7th ed.). Pearson.
- Taber, K. S. (2018). The use of Cronbach's Alpha when developing and reporting research instruments in science education. *Research in Science Education*, 48, 1273-1296. <https://doi.org/10.1007/s11165-016-9602-2>
- Teo, T. (2010). Validation of the technology acceptance measure for pre-service teachers (TAMPST) on a Malaysian sample: A cross-cultural study. *Multicultural Education & Technology Journal*, 4(3), 163-172. <http://doi.org/10.1108/17504971011075165>
- Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). *MIS Quarterly*, 27(3), 425-478. <https://doi.org/10.2307/30036540>
- Vieriu, A. M., & Petrea, G. (2025). The Impact of Artificial Intelligence (AI) on Students' Academic Development. *Education Sciences*, 15(3), 343. <https://doi.org/10.3390/educsci15030343>
- Yim, I. H. Y. (2024). A critical review of teaching and learning artificial intelligence (AI) literacy: Developing an intelligence-based AI literacy framework for primary school education. *Computers and Education: Artificial Intelligence*, 7, 100319. <https://doi.org/10.1016/j.caeai.2024.100319>
- Zhai, C., Wibowo, S., & Li, L. D. (2024). The effects of over-reliance on AI dialogue systems on students' cognitive abilities: A systematic review. *Smart Learning Environments*, 11, 28. <https://doi.org/10.1186/s40561-024-00316-7>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Exploring the Effect of Shape Parameters on Font Design Using Rational Quadratic Trigonometric Curves

Noor Khairiah Razali^{1*}, Nur Ain Fitrah Azhari², Siti Musliha Nor-Al-Din³, Nursyazni Mohd Sukri⁴

^{1,3,4}Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Terengganu Branch, Kuala Terengganu Campus, Malaysia.

²Shell Malaysia Limited, Malaysia.

ARTICLE INFO

Article history:

Received 17 June 2025

Revised 25 August 2025

Accepted 26 August 2025

Online first

Published 1 September 2025

Keywords:

Rational Quadratic
Trigonometric Bézier
Shape Parameter
Font design
Interactive tool

DOI:

10.24191/jcrinn.v10i2.541

ABSTRACT

This study investigates the impact of shape parameters on font contour design using rational quadratic trigonometric Bézier (RQTB) curves. By varying parameters (m , n) and weight (v), the research evaluates contour accuracy through pointwise deviation analysis and computational efficiency. The configuration ($m = 0.5$, $n = -0.5$, $v = 1$) achieved the lowest RMSE and CPU time, indicating optimal performance. An interactive tool was also developed to support real-time parameter adjustment. The results demonstrate that carefully tuned shape parameters enhance both visual fidelity and efficiency, making RQTB curves practical for typography and related design applications.

1. INTRODUCTION

Curve modeling is a core aspect of computer-aided geometric design (CAGD), especially in typography where both visual appeal and precision are required. Rational Bézier curves are widely used because they allow smooth transitions and flexible control through the use of weights (Ceylan et al., 2021; Salomon, 2006). However, one of their main limitations is the lack of local shape control, which makes it difficult to fine-tune detailed parts of designs such as font contours. To address this, rational quadratic trigonometric Bézier (RQTB) curves were introduced. These curves use trigonometric basis functions along with shape parameters, giving designers the ability to adjust curvature directly without having to modify control points (Bashir et al., 2012). This makes them particularly useful in type design, where small changes in curve behavior can have a big impact on readability and style. Previous studies have shown how RQTB curves can be applied in practice. Ahmad and Gobithaasan (2018) introduced RQT spirals that maintain curvature

^{1*} Corresponding author. E-mail address: noorkhairiah@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.541>

continuity for font strokes, while Dube and Gupta (2022) explored how shape parameters affect the appearance of letters. Ahmed et al. (2022) applied Bézier-like curves in Arabic typography, showing that the method is adaptable across different scripts. From a geometric perspective, Sánchez-Reyes (2022) proposed tools for local shape control using shape factors, and Omar et al. (2023) extended the framework with two shape parameters for even greater flexibility. Beyond typography, Munir et al. (2024) applied similar approaches to natural shape modeling, demonstrating the wider relevance of these methods.

Despite these advances, curve modeling still faces important limitations. A common approach to achieving complex shapes is to increase the number of control points. While this can improve approximation, it often results in higher computational cost, more complicated editing, and sometimes unwanted distortions in the outline. In typography, where balance and readability are critical, too many control points can actually reduce the quality of the design. Shape parameters offer a more efficient alternative by allowing designers to adjust curvature directly, without relying on additional control points. This approach not only simplifies the modeling process but also improves accuracy. However, the effectiveness of shape parameters depends on finding the best parameter settings, which has not been studied in a systematic way.

This paper addresses this gap by evaluating how shape parameters (m , n) and weight (v) influence font contour quality and CPU execution time using RQTB curves. A stylized font character was modelled using multiple curve segments, and the accuracy of each configuration was assessed through pointwise deviation analysis, which measures the Euclidean distance between the generated and reference contours. This quantitative approach provides a robust metric for evaluating contour fidelity. Additionally, an interactive tool was developed to allow real-time parameter adjustment, offering immediate feedback on both visual and computational outcomes. The findings demonstrate that carefully tuned shape parameters can significantly enhance both the accuracy and performance of font design, with broader implications for CAD and computer graphics applications.

2. METHODOLOGY

This study applies rational quadratic trigonometric Bézier (RQTB) functions to model font contours, focusing on the influence of global shape parameters on curve geometry. The methodology involves three main stages: constructing curves using predefined control points, adjusting the shape parameters m , n , v and refining the contours through iterative evaluation. Each configuration is assessed based on curvature smoothness, computational efficiency, and visual quality, with contour accuracy measured using pointwise deviation analysis to quantify the alignment between generated curves and the reference font outline.

2.1 Curve formulation

In geometric modeling, rational trigonometric curves offer a powerful framework for representing smooth and flexible shapes, particularly useful in applications like font design where stroke aesthetics are critical. The RQTB is foundational formulations that enable fine control over curve shape through the incorporation of global shape parameters. These basis functions utilize sine and cosine terms to construct curves that can closely mimic circular arcs while also offering designer-defined control over curvature transitions. The RQTB approach is computationally simpler and suitable for applications requiring moderate flexibility.

In this section, the basis sets are defined, and the respective curve formulations are constructed using a set of control points and shape parameters.

2.1.1 Definition 1: Rational Quadratic Trigonometric Bézier (RQTB) basis

According to Bashir et al. (2012), quadratic trigonometric basis functions with two shape parameters, which are m and n , where $m, n \in [-1, 1]$ are defined as:

$$\begin{aligned} f_0(u) &= \left(1 - \sin \frac{\pi}{2} u\right) \left(1 - m \sin \frac{\pi}{2} u\right) \\ f_1(u) &= 1 - \left(1 - \sin \frac{\pi}{2} u\right) \left(1 - m \sin \frac{\pi}{2} u\right) - \left(1 - \cos \frac{\pi}{2} u\right) \left(1 - n \cos \frac{\pi}{2} u\right) \\ f_2(u) &= \left(1 - \cos \frac{\pi}{2} u\right) \left(1 - n \cos \frac{\pi}{2} u\right) \end{aligned} \quad (1)$$

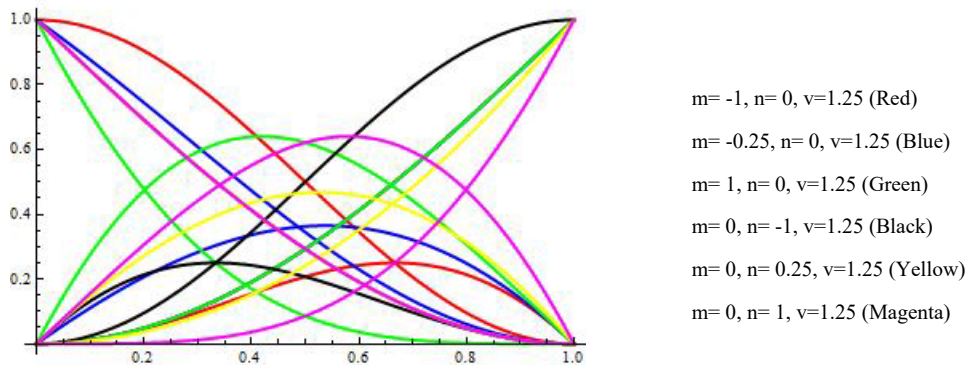


Fig 1: The Basis Graph of RQTB

2.1.2 Definition 2: Rational Quadratic Trigonometric Bézier (RQTB) curve

The RQTB Curve is describing as follows:

$$f(u) = \frac{f_0 P_0 + f_1 P_1 v + f_2 P_2}{f_0 + f_1 v + f_2}, \quad u \in [0, 1], \quad m, n \in [-1, 1] \quad (2)$$

where P_i ($i = 0, 1, 2$) are the control points, f_i ($i = 0, 1, 2$) are the basis functions and v ($v \geq 0$) is the scalar, which also known as the weight of the function (Bashir et al., 2012).

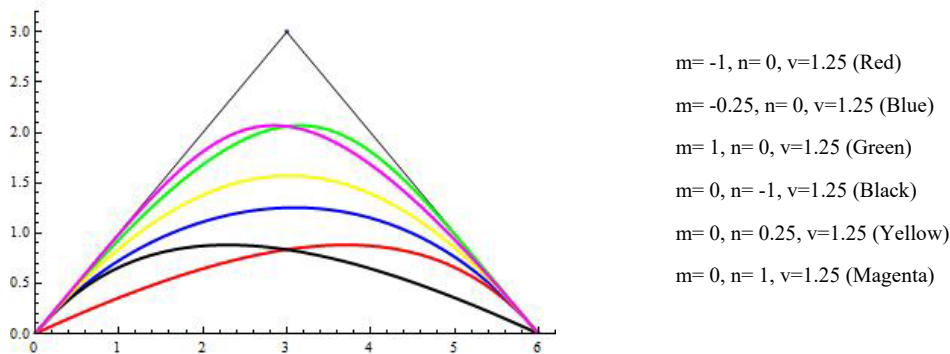


Fig. 2. The curves of RQTB

2.2 The rational trigonometric curve in font design

A stylized script character “V” was chosen as the representative test case for this study. This character was selected because it contains both curved and angular structures, making it a challenging yet suitable subject for testing the flexibility of RQTB curves. Its geometry includes loops, smooth transitions, and sharp inflections features that allow a meaningful evaluation of how global shape parameters affect contour accuracy and smoothness.

The contour was modelled using 20 quadratic rational trigonometric curve segments and 40 control points, as listed in Table 1. These control points were carefully derived from the reference font image (Fig. 3), which served as the benchmark for comparison. Each control point was selected by tracing key structural features of the character, such as curve inflections, stroke tapering, and junctions between different contour regions. The distribution of control points aimed to balance structural fidelity which is accurately capturing the font’s geometry with computational feasibility.

By relying on this structured control point layout, the modelling process ensured that the RQTB curves captured the essential stylistic features of the letter while allowing systematic parameter variation in later analysis.



Fig 3. The real font design (reference figure)

Table 1. The control points of reference figure

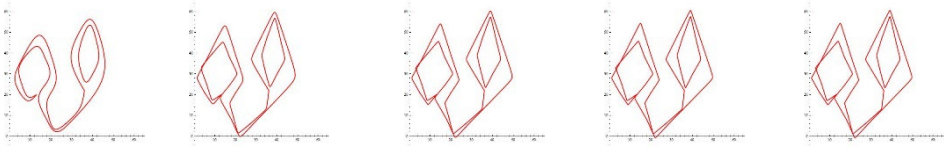
Curve	Control Points P(x, y)	Curve	Control Points P(x, y)
Curve C1	p0 = (13, 20) ; p1 = (9, 15) ; p2 = (5, 21)	Curve C11	p20 = (29.5, 7) ; p21 = (22, -1) ; p22 = (18, 8)
Curve C2	p2 = (5, 21) ; p3 = (1, 28) ; p4 = (6, 37)	Curve C12	p22 = (18, 8) ; p23 = (12, 18.5) ; p24 = (17, 24.5)
Curve C3	p4 = (6, 37) ; p5 = (15, 55) ; p6 = (20, 39)	Curve C13	p24 = (17, 24.5) ; p25 = (21, 30) ; p26 = (17, 38)
Curve C4	p6 = (20, 39) ; p7 = (24, 27) ; p8 = (20, 21)	Curve C14	p26 = (17, 38) ; p27 = (14, 46) ; p28 = (8.5, 38.8)
Curve C5	p8 = (20, 21) ; p9 = (17, 16) ; p10 = (19, 9)	Curve C15	p28 = (8.5, 38.8) ; p29 = (1, 35) ; p30 = (6, 25)
Curve C6	p10 = (19, 9) ; p11 = (22, -1) ; p12 = (27, 6)	Curve C16	p30 = (6, 25) ; p31 = (8.5, 17) ; p32 = (13, 20)
Curve C7	p12 = (27, 6) ; p13 = (35, 13) ; p14 = (36, 22)	Curve C17	p32 = (35, 47) ; p33 = (39, 58) ; p34 = (42, 45)
Curve C8	p14 = (36, 22) ; p15 = (27, 37) ; p16 = (32.5, 48)	Curve C18	p34 = (42, 45) ; p35 = (44, 37) ; p36 = (40, 30)
Curve C9	p16 = (32.5, 48) ; p17 = (39, 61) ; p18 = (43, 49)	Curve C19	p36 = (40, 30) ; p37 = (36.5, 23) ; p38 = (34, 32.5)
Curve C10	p18 = (43, 49) ; p19 = (50, 28) ; p20 = (29.5, 7)	Curve C20	p38 = (34, 32.5) ; p39 = (32.5, 39) ; p40 = (35, 47)

2.3 The influence shape parameter on font design

The shape parameters play a crucial role in rational trigonometric curve modelling by providing designers with intuitive and continuous control over the overall geometry of a curve without modifying the control points. Unlike local shape parameters, which influence only a segment of the curve, global shape parameters affect the entire curve uniformly, making them particularly effective in applications like font design where consistent visual flow and smoothness are important. These parameters denoted as m , n , and weight, v in rational quadratic trigonometric formulations, modulate the influence of basis functions derived from sine and cosine terms. By adjusting these values, designers can sharpen, flatten, or soften curve segments to achieve desired stylistic effects in letter contours. This flexibility allows for the generation of font shapes with varying curvature profiles, making global shape parameters a valuable design tool in balancing artistic style and geometric accuracy.

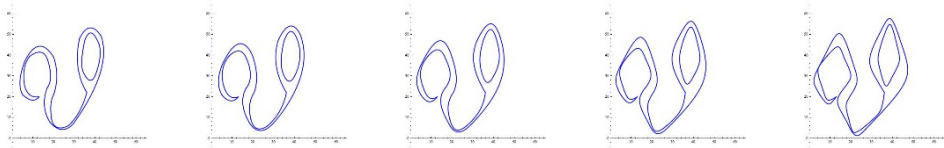
To validate the effectiveness of the generated font contour designs, a visual comparison was conducted against an actual styled font image representing the letter “V” (see Fig. 3). The reference font displays a graceful, fluid structure with a smooth rounded hook on the left and an elegant vertical curve on the right. This image serves as a valuable benchmark for assessing the degree of visual similarity between the generated and actual contours.

Table 2. The font contour for $m = 0.5$, $n = 0.5$ and various values of v

Fonts Contour Designs					
	v	1	5	10	15

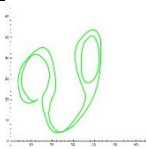
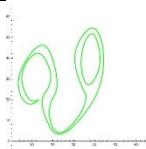
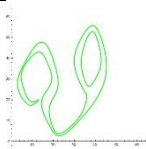
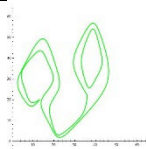
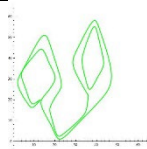
In Table 2, which illustrates the effect of varying the parameter v while holding $m = 0.5$ and $n = 0.5$ constant, the contour produced at $v = 1$ most closely resembles the original font. It retains the desired softness and fluidity in its structure. As the value of v increases, the contour gradually becomes more angular and rigid. Consequently, the visual appeal of the design diminishes, indicating that higher v values are less effective in replicating the original font shape.

Table 3. The font contour for $m = 0.5$, $v = 1$ and various values of n

Fonts Contour Designs					
	n	-1	-0.5	0	0.5

Furthermore, Table 3 demonstrates the influence of parameter n , with $m = 0.5$ and $v = 1$ fixed. Notably, the contours generated using $n = -0.5$ and $n = -1$ exhibit a more refined and proportionate structure. These values introduce subtle inward curvature that enhances the resemblance to the original font. Conversely, when n increases to 0.5 and 1 , the contours appear unbalanced and structurally inconsistent, particularly near the upper loop and lower tail. This indicates that negative values of n are more appropriate for preserving symmetry and design integrity.

Table 4. The font contour for $n = 0.5$, $v = 1$ and various values of m

Fonts Contour Designs					
m	-1	-0.5	0	0.5	1

In addition, Table 4 highlights the effect of varying parameter m while maintaining $n = 0.5$ and $v = 1$. The designs corresponding to $m = -1$ and $m = -0.5$ generate more vertically compact and well-proportioned contours, which align more closely with the actual font's geometry. However, increasing m to 0.5 and 1 leads to vertically stretched and distorted shapes, resulting in diminished visual quality. The loop region becomes particularly exaggerated, detracting from the overall balance and legibility.

In summary, the most accurate and aesthetically consistent results are obtained when $v = 1$, $n = -0.5$ or -1 , and $m = -0.5$ or -1 . These parameter combinations effectively preserve the stylistic features, smooth transitions, and structural harmony characteristic of the actual font. On the other hand, higher values of m and n tend to compromise these qualities, producing uneven contours that deviate from the intended design.

2.4 Pointwise deviation analysis for shape parameter evaluation

To determine the most effective combination of shape parameters and weights in constructing font contours with RQTB, this study applies a pointwise distance analysis between the generated curves and the reference contours. This technique quantitatively evaluates the fidelity of the approximation by measuring how closely each sampled point on the designed curve aligns with the actual outline.

Given a parametric curve constructed from RQTB, denoted by

$$f(u) = (x(u)_c, y(u)_c), \quad u \in [0,1] \quad (3)$$

and the actual contour curve, represented as

$$A(u) = (x(u)_a, y(u)_a) \quad (4)$$

the pointwise deviation is expressed as the Euclidean distance:

$$D_i = \sqrt{(x(u)_c - x(u)_a)^2 + (y(u)_c - y(u)_a)^2} \quad (5)$$

The results are displayed as follows, showing the computed pointwise deviations ($D1$, $D2$, $D3$), along with the Mean and RMSE values for different parameter settings.

Table 5. Pointwise deviation and error analysis for different values of v with $m = n = 0.5$

v	D1	D2	D3	Mean	RMSE
1	3.490702	1.096141	1.900421	2.162421	1.470517
5	6.820359	3.134599	5.760139	5.238366	2.288748
10	7.030348	3.241886	6.080132	5.450789	2.334692
15	8.000306	3.241886	6.080132	5.774108	2.402937
20	7.350333	3.671241	6.620121	5.880565	2.424988

When fixing $m = n = 0.5$ and varying v , the RMSE values range from 1.47 to 2.42. The smallest error is observed at $v = 1$ (RMSE = 1.470517), indicating that the curve generated under this setting best matches the reference contour. As v increases from 5 to 20, the RMSE consistently rises, suggesting that larger values of v reduce the approximation fidelity and lead to a greater deviation from the target outline.

Table 6. Pointwise deviation and error analysis for different values of n with $v = 1$, $m = 0.5$

n	D1	D2	D3	Mean	RMSE
-1	0.912688	1.160078	1.220656	1.097807	1.047763
-0.5	0.485077	0.838492	0.056569	0.460046	0.678267
0	1.23199	0.138351	0.711126	0.693822	0.83296
0.5	3.270749	1.203039	2.22036	2.231383	1.493781
1	5.10048	1.846677	3.610222	3.519126	1.875933

For the case where $v = 1$, $m = 0.5$ and n varies, the results demonstrate that negative values of n provide better approximation accuracy. Specifically, the lowest RMSE occurs at $n = -0.5$ (RMSE = 0.678267), which highlights the effectiveness of slight negative adjustments in improving the curve's alignment. In contrast, positive n values, such as $n = 0.5$ or $n = 1$, show considerably larger RMSE values (1.493781 and 1.875933, respectively), indicating poorer approximation.

Table 7. Pointwise deviation and error analysis for different values of m with $v = 1$, $n = 0.5$

m	D1	D2	D3	Mean	RMSE
-1	0.703491	0.946322	0.900888	0.850234	0.922081
-0.5	0.783135	0.838492	0.056569	0.559398	0.747929
0	1.991231	0.054231	1.140702	1.062054	1.03056
0.5	3.490702	1.905656	2.320345	2.572234	1.603819
1	4.990491	2.061406	3.400235	3.484044	1.866559

Similarly, for $v = 1$, $n = 0.5$ with varying m , the accuracy trend is consistent. The best performance is achieved at $m = -0.5$ (RMSE = 0.747929), followed by $m = -1$ (RMSE = 0.922081). Positive values of m , particularly $m = 0.5$ and $m = 1$, result in higher RMSE values (1.603819 and 1.866559), indicating reduced fidelity of the constructed curve.

Overall, the findings indicate that lower or slightly negative values of the parameters m and n , along with smaller values of v , are more effective in minimizing pointwise deviations. This suggests that fine-tuning the shape parameters towards negative or near-zero ranges contributes to better alignment of the constructed curves with the actual contours, enhancing the precision of the RQTB method in font contour design.

2.5 The computation time on font design

This study focuses on evaluating the computational efficiency of RQTB methods in font curve design by analysing CPU time with varying global shape parameters (m , n) and weight (v). By systematically adjusting these parameters, the study examines how changes in curve flexibility and stroke refinement impact the time required for curve generation and curvature computation. The CPU time metric serves as a secondary aspect for assessing the computational cost associated with different shape parameter settings.

This focused analysis aims to identify which parameter configurations optimize performance, offering insights into the trade-offs between shape control and computational efficiency in rational trigonometric curve formulations.

Table 7. The Computation Time of the Image

Shape Parameter, v	CPU Time	Shape Parameter, n	CPU Time	Shape Parameter, m	CPU Time
1	0.07	-1	0.069	-1	0.085
5	0.115	-0.5	0.068	-0.5	0.088
10	0.111	0	0.063	0	0.068
15	0.102	0.5	0.069	0.5	0.091
20	0.103	1	0.069	1	0.088
Average CPU Time	0.1002	Average CPU Time	0.0676	Average CPU Time	0.084

Table 7 shows that the choice of shape parameters significantly affects both CPU time and contour quality. For parameter v, the best contour with the lowest computation time was achieved at $v = 1$, compared to the overall average of 0.1002 seconds. For parameter n, although the minimum CPU time occurred at $n = 0$, the best contours were observed at $n = -0.5$ and -1 , both close to the average of 0.0676 seconds. Similarly, for parameter m, values of -0.5 and -1 yielded better contours than other settings, with CPU times slightly above the average of 0.084 seconds. In conclusion, the shape parameter of $v = 1$, $n = -0.5$, and $m = -0.5$ provides the best trade-off between high-quality contours and efficient computation.

2.6 The interactive tool for shape parameter

To facilitate a deeper understanding of the influence of the shape parameters on rational trigonometric font curves, an interactive tool has been developed. This tool allows users to dynamically adjust shape parameters and observe the immediate effects on the resulting curves. By incorporating this tool, users can experiment with different values for the shape parameters, specifically m, n, and weight, v, and see how these changes impact the smoothness, curvature, and overall appearance of font contours. Additionally, the tool provides real-time feedback, including curvature plots and CPU time measurements, offering valuable insights into the trade-offs between computational efficiency and visual aesthetics in font design.

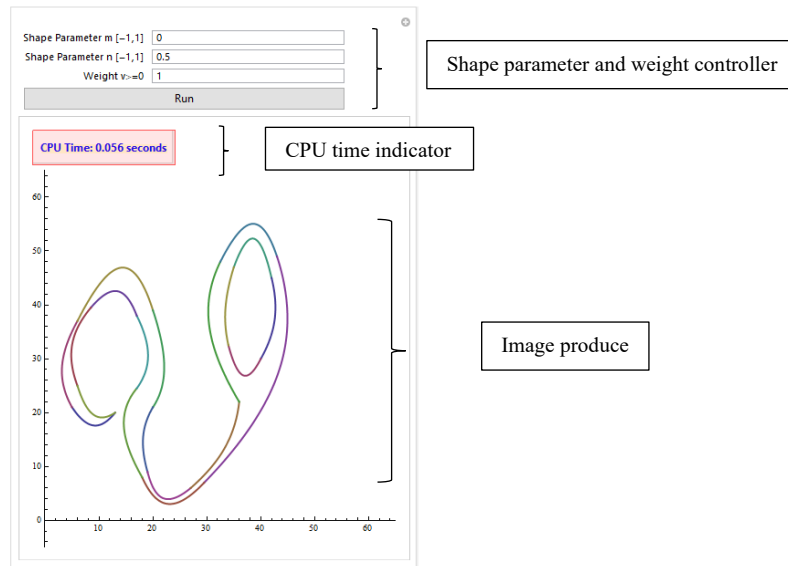


Fig 4. The interactive tool for shape parameter

3. RESULTS AND DISCUSSION

In this study, a selected set of shape parameter values was tested to reconstruct the font character using 20 RQTB segments and 40 control points. Among the tested configurations, a particular set of parameters produced the closest approximation to the actual font shape with smooth curvature and visually consistent strokes. This design also recorded the lowest CPU time, indicating efficient computation. Although not all possible parameter values were explored, the results show that high-quality font curves can be achieved with minimal adjustment, demonstrating the practicality and effectiveness of the RQTB method for font design.

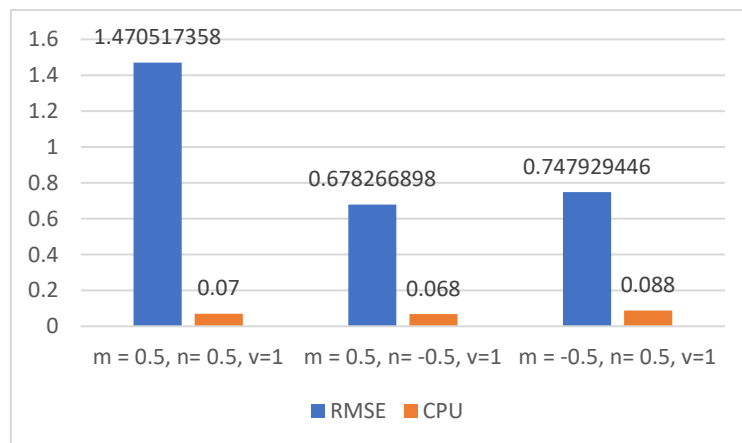
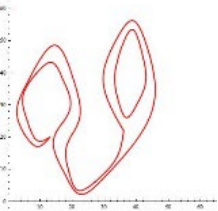
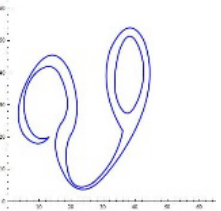
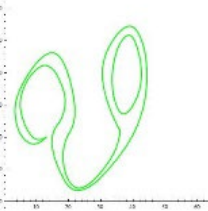


Fig 5. Evaluation of shape parameters based on RMSE and CPU time

Table 8. The result of font contour design

IMAGE				
SHAPE PARAMETER	$M = 0.5, N = 0.5, V = 1$	$M = 0.5, N = -0.5, V = 1$	$M = -0.5, N = 0.5, V = 1$	

The analysis of shape parameter combinations reveals a clear trade-off between contour accuracy and computational efficiency. Among the tested configurations, the combination $m = 0.5$, $n = -0.5$, $v = 1$ achieved the lowest RMSE value (0.678), indicating the highest fidelity to the original font contour. This configuration also recorded the lowest CPU time (0.068s), making it the most efficient and accurate overall.

Similarly, the combination $m = -0.5$, $n = 0.5$, $v = 1$ produced a smooth and visually consistent contour with a slightly higher RMSE (0.748) and CPU time (0.088s), still within acceptable performance limits. In contrast, configurations such as $m = 0.5$, $n = 0.5$, $v = 1$ resulted in significantly higher RMSE (1.4705), indicating poor contour approximation despite a relatively low CPU time (0.07s). The combined RMSE and CPU time graph clearly illustrates that lower or slightly negative values of m and n , paired with $v = 1$, consistently yield better performance in terms of both accuracy and speed. These findings support the use of RQTB curves with optimized shape parameters for efficient and high-quality font design.

4. CONCLUSION

This study investigated the impact of shape parameters m , n and weight, v on font contour design using RQTB curves. Through systematic evaluation of contour fidelity and computational efficiency, the findings demonstrate that optimal font reconstruction is achieved when shape parameters are carefully tuned. Specifically, the combinations $m = 0.5$, $n = -0.5$, $v = 1$ and $m = -0.5$, $n = 0.5$, $v = 1$ produced the most accurate contours with minimal deviation from the original font shape, while also maintaining low CPU time. The comparative analysis, supported by RMSE and CPU time graphs, highlights a clear trade-off between accuracy and performance. Lower or slightly negative values of m and n , paired with $v = 1$, consistently yielded superior results. These configurations not only preserved the stylistic integrity of the font but also ensured efficient computation, making them ideal for practical applications in digital typography, CAD, and computer graphics. Future work may explore adaptive parameter selection techniques, integration with machine learning models for automated tuning, and the extension of RQTB curves to higher-degree or fractional forms to enhance flexibility and visual fidelity in complex font designs.

5. ACKNOWLEDGEMENTS/FUNDING

The authors would like to acknowledge Universiti Teknologi MARA Cawangan Terengganu Kampus Kuala Terengganu for their support.

6. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

7. AUTHORS' CONTRIBUTIONS

Noor Khairiah Razali: Conceptualisation, methodology, formal analysis, investigation and writing-original draft; **Nur Ain Fitrah Azhari:** Conceptualisation, methodology, and formal analysis; **Siti Musliha Nor-Al-Din:** Conceptualisation, formal analysis, and validation; **Nursyazni Mohd Sukri:** Conceptualisation, supervision, writing- review and editing, and validation.

8. REFERENCES

- Ahmad, A., & Gobithaasan, R. (2018). Rational quadratic Bézier spirals. *Sains Malaysiana*, 47(9), 2205–2211. <https://doi.org/10.17576/jsm-2018-4709-31>
- Ahmed, I. A. A., Omar, A. H. A., & Ali, J. M. (2022). Arabic font design using quadratic Bézier-like curve. *Journal of Pure & Applied Sciences*, 21(4), 50–56.
- Bashir, U., Abbas, M., Abd Majid, A., & Ali, J. M. (2012, July). The rational quadratic trigonometric Bézier curve with two shape parameters. In *2012 Ninth International Conference on Computer Graphics, Imaging and Visualization* (pp. 31–36). IEEE. <https://doi.org/10.1109/CGIV.2012.21>
- Ceylan, A. Y., Turhan, T., & Tükel, G. O. (2021). On the geometry of rational Bézier curves. *Honam Mathematical Journal*, 43(1), 88–99. <https://doi.org/10.5831/HMJ.2021.43.1.88>
- Dube, M., & Gupta, N. (2022). Font design through RQT Bézier curve using shape parameters. In *Advances in Mathematical Modelling, Applied Analysis and Computation*. In *Proceedings of ICMMAAC 2021* (pp. 319–335). Springer. https://doi.org/10.1007/978-981-16-8641-1_23
- Munir, N. A. A. A., Hadi, N. A., Nasir, M. A. S., & Halim, S. A. (2024, January). Modelling fruit outline using cubic trigonometric spline. In *AIP Conference Proceedings* (Vol. 2905, No. 1). AIP Publishing. <https://doi.org/10.1063/5.0172468>
- Omar, A., Ahmed, I., & Abuzaiyan, F. (2023). The quadratic trigonometric Bézier-like curve with two shape parameters. *Journal of Pure & Applied Sciences*, 22(3), 215–221.
- Salomon, D. (2006). *Curves and surfaces for computer graphics*. Springer. <https://doi.org/10.1007/0-387-28452-4>
- Sánchez-Reyes, J. (2022). The uniqueness of the rational Bézier polygon. *Computer Aided Geometric Design*, 96, 102118. <https://doi.org/10.1016/j.cagd.2023.102266>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Key Determinants of Athletic Success: A Fuzzy AHP Approach

Khairu Azlan Abd Aziz¹, Wan Suhana Wan Daud^{2*}, Mohd Fazril Izhar Mohd Idris³, Muhammad Amsyar Izmi⁴, Rizauddin Saian⁵

^{1,3,4,5}*Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Perlis Branch, Arau Campus, 02600 Arau Perlis, Malaysia.*

²*Department of Mathematical Sciences, Faculty of Intelligent Computing, Universiti Malaysia Perlis, Kampus Alam UniMAP, Pauh Putra, 02600 Arau, Perlis, Malaysia.*

ARTICLE INFO

Article history:

Received 28 March 2025

Revised 26 August 2025

Accepted 27 August 2025

Published 1 September 2025

Keywords:

Athletic Success

Saaty's Scale

Fuzzy Analytic Hierarchy Process

DOI:

10.24191/jcrinn.v10i2

ABSTRACT

Athletic success is influenced by multiple factors, requiring a structured approach to identify and prioritize the key contributors to success. The Fuzzy Analytical Hierarchy Process (AHP) method provides a systematic decision-making framework to evaluate these factors by incorporating expert judgment and handling uncertainty in pairwise comparisons. This study employs the fuzzy AHP method to assess and rank the factors influencing an athlete's successful in sports. Two experts, both lecturers from the Faculty of Sports Science and Recreation at UiTM Perlis, were involved in the evaluation due to their extensive experience and achievements in sports, making them suitable decision-makers for this study. Based on the findings, the factor of facilities plays a crucial role in sports performance, with a normalized weight of 0.4658, while coaching is identified as the most critical sub-factor, receiving a normalized weight of 0.2399. These results demonstrate that Fuzzy AHP effectively handles uncertainty and subjective judgments, providing valuable insights for athletes, coaches, and sports management in optimizing training and resource allocation.

1. INTRODUCTION

An athlete is a person who is trained or skilled in sports, physical activities, or competitive games that require strength, agility, endurance, or coordination. Athletes can participate in various sports, ranging from individual events like running and swimming to team sports like football and basketball. They undergo rigorous training to improve their performance and often compete at different levels, from amateur to professional and even international competitions like the Olympics.

Achieving success as an athlete requires a combination of various factors, including natural talent, physical fitness, technical skills, discipline, access to sports facilities and psychological strength. These

^{2*} Corresponding author. *E-mail address:* wsuhana@unimap.edu.my
<https://doi.org/10.24191/jcrinn.v10i2>

factors may influence the overall performance of an athlete and determine their ability to compete at higher levels (Kubiak, 2012). However, prioritizing these factors can be complex, as they often involve subjective judgments and uncertainties. This is because each type of sport requires a different level of emphasis on each factor, whereby each athlete may respond differently to these factors. Additionally, factors such as psychological strength and talent are difficult to measure objectively, making the decision-making process more challenging. The influence of external elements, particularly access to training facilities, further adds to the complexity of analysing the most important factors for athletic success.

Once complexity and ambiguity situation occur in the process of analysing or evaluation, one of the most common approaches will be applied is the fuzzy Analytic Hierarchy Process (Fuzzy AHP). This method was integrated based on the classical Analytic Hierarchy Process (AHP) by Van Laarhoven and Pedrycz (1983). They extended the traditional AHP by incorporating fuzzy logic to handle uncertainty and imprecision in decision-making. Their approach used fuzzy numbers to represent pairwise comparisons, making AHP more suitable for real-world problems with vague or subjective judgments.

Specifically, this study aims to identify the key factors and sub-factors that contribute to athletic success, rank the main factors using the fuzzy AHP method, and further rank the sub-factors accordingly. This study focuses on four primary factors that athletes should consider in achieving success, which are discipline, facilities, talent, and psychology. The analysis will be conducted based on the sub-factors associated with each of these main factors. It is hoped that this study will provide valuable insights for athletes by highlighting the most critical factors for success. Additionally, it will offer significant contributions to the Ministry of Youth and Sports by providing a structured analysis to guide sports development programs and talent management strategies.

This paper is organized as follows. In Section 2, some literature reviews are provided on the fuzzy AHP and the factors that related to athletic success. The methodology used for the study is provided in Section 3. Next, the results and discussion are presented in Section 4 and finally in Section 5, the conclusion is drawn.

2. LITERATURE REVIEW

This section outlines the theories behind fuzzy AHP and the factors considered in this study.

2.1 Fuzzy analytical hierarchy process

The Fuzzy AHP is an extension of the traditional AHP that incorporates fuzzy logic to handle uncertainty and imprecision in decision-making. Basically, the fuzzy AHP method relies on data collected through questionnaires or interviews, typically from experts with relevant knowledge in the field. These experts provide valuable insights that shape the decision-making process. As Saaty and Ozdemir (2015) highlight, fuzzy AHP is based more on expert opinions than on strict statistical analysis. This means that the number of experts involved matters less than their level of expertise in the subject.

One of the strengths of fuzzy AHP is its ability to handle pairwise comparisons, which help structure complex decisions by ranking different factors in a hierarchical way (Zabihi et al., 2020). Since AHP captures human judgment, which is often complicated to handle, it has frequently been selected to be applied in solving the multi-criteria decision-making (MCDM) problems.

As for now, researchers have applied fuzzy AHP in many areas. Some examples include selecting contractors (Rashid et al., 2020), identifying groundwater potential zones in Bangladesh (Sresto et al., 2021), assessing land conflict risks (Peng et al., 2021), choosing effective strategies for COVID-19 prevention (Idris et al., 2023), evaluating nasyid competitions (Aziz et al., 2023), and selecting the best student award (Aziz et al., 2023). More recently, in 2024, there is a study that has explored how Fuzzy SERVQUAL methods can be integrated into Fuzzy AHP to measure service quality and satisfaction in

educational programs (Aulia et al., 2024). In summary, based on the significant number of studies, Fuzzy AHP has proven to be consistently relevant in real-world applications.

2.2 Factors influencing athletes' success

There are many factors that can influence an athlete's success. In this study, four key factors are considered, which are the athlete's discipline, natural talent, available facilities, and psychological aspects. A literature review has been conducted to explore these factors.

2.2.1 Discipline

According to Denison et al. (2017), discipline is one of the most important criteria that athletes should consider. Discipline is essential in any sport because it builds an athlete's attitude, helping them stay focused and achieve their goals without distractions. It also provides athletes with a set of rules that they can implement in their daily lives, allowing them to live efficiently and effectively. When athletes develop discipline, they can make small sacrifices for a better future. One crucial aspect of discipline is understanding and adhering to rules. Rules provide structure and guidelines that help athletes stay focused on their goals (Graham & Burns, 2020). By following the rules of their sport, athletes learn to channel their efforts toward specific objectives.

Another crucial factor in an athlete's discipline is good time management (Macquet & Skalej, 2015). Athletes who manage their time well can balance training, competitions, and personal life more effectively. They need to plan their daily schedule by setting time for training, rest, proper meals, and recovery sessions. However, they often face challenges such as a tight training schedule and academic responsibilities. Therefore, using a schedule or a to-do list can help them organize their activities better. Good time management not only improves performance in sports but also helps build a more responsible character.

2.2.2 Facilities

The important facilities for athletes include training venues, sports equipment, and quality coaches. A suitable venue with basic facilities such as fields, courts, tracks, or swimming pools that meet international standards is crucial to ensure effective training (Nacar et al., 2013). In addition, extra facilities like changing rooms, rest areas, and physiotherapy spaces are also needed to maintain the comfort and well-being of athletes. A good venue not only provides a comfortable training space but also helps improve an athlete's performance in competitions by offering a safe and conducive environment for practice and competition at both national and international levels.

Besides that, sports equipment is also very important in helping an athlete succeed (Diejomaoh et al., 2015). High-quality and suitable equipment for each sport can improve performance and reduce the risk of injuries. For example, running shoes designed for better support help protect the feet, while lightweight and durable badminton rackets allow players to control their movements more effectively. Regular maintenance of equipment is also essential to ensure it remains in good condition and safe to use.

A good coach is also an important element in an athlete's success. Coaches help athletes improve their skills, techniques, and performance by providing technical knowledge, strategic guidance, and structured training. They not only understand the technical aspects of the sport but can also adjust training sessions according to the athlete's level and needs (Suyudi, 2023). In addition, experienced coaches can provide motivation, develop discipline, and guide athletes in both mental and physical aspects to ensure they are always in their best condition for training and competitions.

2.2.3 Talent

Talent in sport refers to an individual's natural ability or potential to excel in a specific sport, which can be seen in their technical skills and physical attributes (such as strength, speed, and endurance). It is often a combination of genetic traits, such as strength, speed, and agility, along with mental attributes like determination, focus, and resilience. However, talent alone is not enough for success, it must be developed through consistent training, proper coaching, and a supportive environment. According to Wilson et al. (2017), players with natural talent, combined with greater skill and balance, were more likely to perform better when given proper training. Additionally, athletes with great talent often show a quick ability to learn new techniques, adapt to different game situations, and perform at a high level under pressure.

2.2.4. Psychology

Psychology helps athletes set meaningful goals and maintain motivation throughout their training and competition journey (Bulent et al., 2017). Every athlete needs strong psychological well-being in their life. There are three key psychological factors that athletes should consider, with the first being motivation. Ibrahim et al. (2016) states that motivation provides athletes with the determination and focus needed to achieve their goals. Whether aiming to win a championship, break a personal record, or qualify for the Olympics, motivation keeps athletes focused on their objectives. It helps them maintain a clear vision of what they want to accomplish and drives their efforts toward success.

The second factor is mental rehabilitation, which involves mental training to prepare an athlete's mind for peak performance, both mentally and physically. Covassin et al. (2015) highlights that mental attributes such as confidence, focus, self-belief, and motivation are crucial for success in sports. These factors can help athletes reach higher levels of performance when combined with physical talent. Lastly, goal setting plays a crucial role in providing athletes with clear direction and focus. Baker et al. (2003) explains that setting specific goals, such as improving performance, winning an award, or reaching a particular ranking, help athletes channel their efforts, prioritize training, and align their actions with their objectives.

3. RESEARCH METHODOLOGY

The methodology of the study is designed as illustrated in the following Fig. 1.

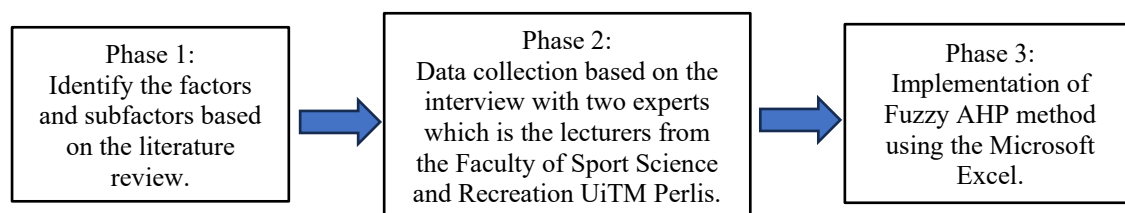


Fig. 1. Research methodology phases

The detail description is given as follows.

3.1 Phase 1: Identify the criteria and sub-criteria based on the literature review

As discussed in Subsection 2.2, the main factors (criteria) and their corresponding sub-factors (sub-criteria) influencing athletes' success were identified from past studies. These elements, which have been elaborated in the literature review, are summarized in Table 1 for clarity.

Table 1. Factors and sub-factors of athletes' success

Factor	Sub-factors
Discipline	Time -management
	Attitude
	Rules
Facilities	Equipment
	Sport facilities
	Coaching
Talent	Skill
	Fitness
	Physical
Psychology	Motivation
	Mental rehab
	Goal setting

3.2 Phase 2: Data collection

The data collection process involves conducting interviews with experts in the field. In this study, the experts consist of two lecturers from the Faculty of Sport Science and Recreation, UiTM Perlis, who have extensive experience and notable achievements in sports. These experts provide valuable insights based on their expertise in athletic performance, coaching, and sports management. The interviews were conducted using a structured questionnaire that consisted of two sections: a demographic section (collecting information such as name, position, and qualification) and the main questionnaire. Section A contained questions comparing each main factor to determine effective ways of achieving success as an athlete, while Section B involved questions comparing the sub-factors in relation to each main factor.

3.3 Phase 3: Implementation of Fuzzy AHP

In this study, Fuzzy AHP is implemented using Microsoft Excel, which offers a practical and accessible approach for analyzing complex multi-criteria decision-making problems without the need for specialized software. The following section presents a step-by-step guide to the implementation process in Excel.

Step 1: Construct the pairwise comparison matrix (PCM)

In constructing the pairwise comparison matrix (PCM), the Saaty scale (1 to 9) is used, according to the following Table 2.

Table 2. Linguistic variable for pairwise comparison of each criterion

Scale	Reciprocal	Linguistic Variables
1	1	Equally Important
2	1/2	Equally to weakly important
3	1/3	Weakly important
4	1/4	Weakly to fairly important
5	1/5	Strongly Important
6	1/6	Fairly to strongly important
7	1/7	Strongly Important
8	1/8	Strongly to absolutely important
9	1/9	Absolutely Important

In Microsoft Excel, these values are entered in a matrix format using separate column for lower, middle and upper fuzzy numbers, which based on the judgement provided by the experts. The general form of the PCM is given as follows.

$$A = \begin{bmatrix} 1 & a_{12} & \cdots & a_{1n} \\ a_{21} & 1 & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \cdots & 1 \end{bmatrix}$$

where $a_{11}, a_{12}, \dots, a_{nn}$ represent the input or values assigned by the experts based on a scale of 1 to 9. These values correspond to linguistic variables, as shown in the Table 2. It is important to note that $a_{ji} = 1/a_{ij}$ and $a_{ii} = 1$ for every $i, j = 1, 2, 3, \dots, n$. In other words, if the preferences value p_{ij} appears in the upper triangle of the matrix, then its reciprocal $a_{ji} = 1/a_{ij}$ must be in the lower triangle, and vice versa. As the result, the PCM denoted as matrix A , which is always positive and symmetric (Bozanic & Pamucar, 2013).

Step 2: Calculate the largest eigenvalue

Before obtaining the largest eigenvalue denoted as $\lambda_{\max}$, it is necessary to determine the priority vector, which is obtained by averaging the rows of the normalized PCM. Then, the weighted sum values are computed by multiplying the values in the PCM with the priority vector. From that, the values of W are then determined using the following Eq. (1).

$$W = \frac{\text{weighted sum value}}{\text{criteria weight}} \quad (1)$$

Hence, $\lambda_{\max}$ is obtained by averaging the values in W .

Step 3: Compute the Consistency Ratio (CR)

The value of the CR is used to verify the consistency of judgements, where it must be 0.1 or lower (Saaty, 1980). Otherwise, PCM should be revised, or data collection should be reconducted. The formula for CR is determined using Eqs (2) and (3).

$$CR = \frac{\text{Consistency index}(CI)}{\text{Random consistency index}(RI)} \quad (2)$$

where

$$CI = \frac{\lambda_{\max} - n}{n - 1} \quad (3)$$

Here, n denotes the total number of factors. The random consistency index (RI), which depends on the number of factors, is provided in Table 3 (Saaty, 1980).

Table 3. Random consistency index

Number of factors, n	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
Ratio Index, RI	0.00	0.00	0.58	0.90	1.12	1.24	1.32	1.41	1.45	1.49	1.51	1.48	1.56	1.57	1.58

Step 4: Performed the fuzzy pairwise comparison matrix

Next, the entries in the PCM based on the interview are converted into their corresponding triangular fuzzy numbers (TFNs) of (l, m, u) as shown in Table 4.

Table 4. Linguistic variable for the fuzzy pairwise comparison of each criterion

Classical Saaty' Scale	Triangular Fuzzy Number	Triangular Fuzzy Reciprocal Number	Linguistic Variables
1	(1, 1, 1)	(1, 1, 1)	Equally Important
2	(1, 2, 3)	(1/3, 1/2, 1)	Weakly Important
3	(2, 3, 4)	(1/4, 1/3, 1/2)	Fairly Important
4	(3, 4, 5)	(1/5, 1/4, 1/3)	Strongly Important
5	(4, 5, 6)	(1/6, 1/5, 1/4)	Absolutely Important
6	(5, 6, 7)	(1/7, 1/6, 1/5)	Intermittent Values
7	(6, 7, 8)	(1/8, 1/7, 1/6)	
8	(7, 8, 9)	(1/9, 1/8, 1/7)	
9	(9, 9, 9)	(1/9, 1/9, 1/9)	

Thus, the PCM becomes a fuzzy PCM, $\tilde{A} = (l, m, u)$ as follows.

$$\tilde{A} = \begin{bmatrix} (1,1,1) & \tilde{A}_{12} & \cdots & \tilde{A}_{1n} \\ \tilde{A}_{21} & (1,1,1) & \cdots & \tilde{A}_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \tilde{A}_{n1} & \tilde{A}_{n2} & \cdots & (1,1,1) \end{bmatrix}$$

Since there were two experts, each expert's preference had to be averaged and calculated using the Eq. (4).

$$Average(l_{ij}, m_{ij}, u_{ij}) = \frac{\sum_{k=1}^k (l_{nn}^k, m_{nn}^k, u_{nn}^k)}{k} \quad (4)$$

where k represents the number of experts.

Step 5: Calculate the geometric mean

The fuzzy geometric mean $\tilde{r}_i$ for each factor i is determined using the following Eq. (5).

$$\tilde{r}_i = (\tilde{a}_{i1} \times \tilde{a}_{i2} \times \dots \times \tilde{a}_{in})^{\frac{1}{n}} \quad (5)$$

Then, the vector summation of the geometric mean and its reciprocal is computed by employing the following Eqs. (6) and (7) respectively.

$$\text{Vector summation} = \sum_{i=1}^n \tilde{r}_i = \left(\sum l_{r_i}, \sum m_{r_i}, \sum u_{r_i} \right) \quad (6)$$

$$\text{Reciprocal} = \left(\sum_{i=1}^n \tilde{r}_i \right)^{-1} = \left(\frac{1}{\sum u_{r_i}}, \frac{1}{\sum m_{r_i}}, \frac{1}{\sum l_{r_i}} \right) \quad (7)$$

Next, the fuzzy weight of the vector $\tilde{w}_i$ is obtained by multiplying the geometric mean $\tilde{r}_i$ with the reciprocal obtained in Eq. (7) with in ascending order. The formula of the fuzzy weight is given by the Eq. (8).

$$\tilde{w}_i = \tilde{r}_i \times \left(\sum_{i=1}^n \tilde{r}_i \right)^{-1} \quad (8)$$

Then defuzzification of the $\tilde{w}_i$ is calculated using Eq. (9).

$$w_i = \frac{l_i + m_i + u_i}{3} \quad (9)$$

Step 7: Ranking of the factor

The ranking and selection of decisions are conducted based on the normalized non-fuzzy weights. The factors are arranged in descending order, with the highest value representing the most significant factor. The ranking is determined using the following Eq. (10).

$$Z_i = \frac{w_i}{\sum_{i=1}^n w_i} \quad (10)$$

The factor with the highest normalized weight is considered the most important. All steps are repeated to all sub-factors that influence factors for athletes' success. Meanwhile, to calculate the normalization of sub-factor, the weight of sub-factor must be multiplied with the weight of factors to form global weight using the Eq. (11).

$$\text{Global Weight} = \text{Weight of Factor} \times \text{Weight of Sub-factor} \quad (11)$$

4. RESULT AND DISCUSSION

This section presents and discusses the analysis of the factors and sub-factors that contribute to athletes' success. The results are presented in Table 5 as follows.

Table 5. Fuzzy weight, non-fuzzy weight, and normalized weight of all factors

Factor	Fuzzy Weight	Non-Fuzzy Weight	Normalized Weight	Rank
Discipline	(0.1495, 0.2565, 0.4040)	0.2700	0.2539	2
Facilities	(0.3223, 0.4733, 0.6903)	0.4953	0.4658	1
Talent	(0.1444, 0.2242, 0.3738)	0.2475	0.2327	3
Psychology	(0.0327, 0.0460, 0.0729)	0.0505	0.0475	4

According to Table 5, the fuzzy AHP method identifies facilities as the primary factor for making successful athletes, with a normalized weight of 0.4658. The subsequent factor is discipline, with a normalized weight of 0.2539, followed by talent with a normalized weight of 0.2327, and psychology with a normalized weight of 0.0475. The experts agreed that these facilities play a crucial role in enabling athletes to enhance their skills, improve their physical fitness, and practice under optimal conditions. High-quality facilities offer athletes a suitable environment for training. Having safe facilities allows athletes to focus on their performance without unnecessary concern for their well-being. State-of-the-art facilities can provide athletes with an edge by offering cutting-edge technology and equipment. For example, advanced sports science laboratories can assess athlete's physical attributes, monitor their performance, and provide data-driven insights to enhance training methods. Specialized facilities like swimming pools with wave machines or running tracks with force plates can aid in specific training needs.

On the other hand, discipline takes the second spot, emphasizing the vital role of commitment, work ethic, and adherence to training and competition routines. Athletes who maintain a focused and dedicated approach, follow structured regimens, and exhibit self-control and determination are more likely to achieve consistent performance, skill development, and long-term goals. Talent is ranked third, recognizing the significance of natural ability, skill, and aptitude in sports. Athletes with inherent attributes and physical capabilities possess a competitive edge, which serves as a foundation for further growth through training, practice, and skill refinement. Finally, psychology holds the fourth position, highlighting the importance of the mental aspect of sports. Sports psychology helps athletes cultivate a resilient mindset, motivation, focus, and confidence. Mental strength and psychological well-being are crucial for consistent performance, handling pressure, and overcoming challenges.

Subsequently, the rank of the sub-factors is presented in the following Table 6.

Table 6. Fuzzy weight, non-fuzzy weight, and normalized weight of all sub-factors

Factor	Weight of Factor	Sub-factor	Weight of Sub-factor	Weight of Factor × Weight of Sub-factor	Normalized Weight of Sub-factor	Rank
Discipline	0.2539	Time-Management	0.2535	0.0644	0.0644	8
		Attitude	0.2887	0.0733	0.0733	6
		Rules	0.4578	0.1162	0.1163	4
		Equipment	0.2162	0.1007	0.1007	5
Facilities	0.4658	Sport Facilities	0.2689	0.1253	0.1253	3
		Coaching	0.5149	0.2398	0.2399	1
Talent	0.2327	Skill	0.1030	0.0240	0.0240	10
		Fitness	0.2938	0.0684	0.0684	7

Psychology	0.0475	Physical	0.6031	0.1403	0.1404	2
		Motivation	0.3514	0.0167	0.0167	11
		Mental				
		Rehab	0.0728	0.0035	0.0035	12
		Goal setting	0.5757	0.0273	0.0273	9

Table 6 presents the ranking of sub-factors contributing to an athlete's success. Coaching is identified as the most important factor with a normalized weight of 0.2399, followed by physical attributes (0.1404), sport facilities (0.1253), rules (0.31163), equipment (0.1007), attitude (0.0733), fitness (0.0684), time management (0.0644), goal setting (0.0273), skill (0.0240), motivation (0.0167), and mental rehabilitation (0.0035).

Coaching plays a vital role in athlete development, providing guidance, technical expertise, and mentorship to enhance performance. Physical attributes, including endurance, strength, flexibility, and body composition, are essential for optimal performance and injury prevention. Access to well-equipped sport facilities supports effective training and skill development, while adherence to rules ensures fair competition and sportsmanship. Proper equipment is crucial for safety and performance, with sport-specific gear enhancing efficiency and effectiveness.

An athlete's attitude influences their determination, resilience, and discipline, all of which contribute to long-term success. Maintaining fitness is equally important, as it enhances endurance and reduces the risk of injuries. Effective time management allows athletes to balance training, recovery, and personal commitments efficiently. Goal setting provides direction and motivation, ensuring athletes remain focused on their improvement. Skill development involves mastering sport-specific techniques and tactical understanding, which are critical for competition. Motivation sustains an athlete's effort and helps overcome challenges, whether driven by personal fulfillment or external rewards. Lastly, mental rehabilitation supports athletes in managing stress, setbacks, and performance anxiety, fostering mental resilience and overall well-being.

In conclusion, achieving success in sports requires a holistic approach that integrates coaching, physical conditioning, access to facilities, adherence to rules, proper equipment, and mental resilience. While the importance of each sub-factor may vary based on the sport and individual circumstances, their balanced application contributes to overall growth and athletic achievement.

5. CONCLUSION

Based on the objectives stated in the study, the main factors and sub-factors contributing to athlete success were identified and ranked using the fuzzy AHP method. The data for the study were collected from two expert lecturers in the Faculty of Sport Science and Recreation, chosen based on their experience and achievements in sports.

The fuzzy AHP method, which combines fuzzy theory and AHP, effectively achieved the study's objectives by identifying and ranking the main factors and sub-factors that play a crucial role in athlete success. The findings indicate that facilities emerged as the most important factor, followed by discipline, talent, and psychology. Among the sub-factors, coaching exhibited the highest significance, while mental rehab had the lowest level of importance.

In conclusion, this study successfully fulfilled its objectives of identifying and ranking the main factors and sub-factors influencing athlete success using the fuzzy AHP method. The results provide valuable insights for understanding the key elements that contribute to athlete success and can inform decisions related to athlete development and performance.

6. ACKNOWLEDGEMENT

The authors would like to acknowledge the support of Universiti Teknologi Mara (UiTM), Cawangan Perlis and Department of Mathematical Sciences, Faculty of Intelligent Computing, UniMAP for providing the facilities support on this research. The authors would also like to express their gratitude to the anonymous referee for the constructive comments to improve this study.

7. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

8. AUTHORS' CONTRIBUTION

Khairu Azlan Abd Aziz: Supervision, conceptualisation, development of methodology, and formal analysis; **Wan Suhana Wan Daud:** Data interpretation, manuscript writing, and final editing.; **Mohd Fazril Izhar Mohd Idris:** Conceptualisation, methodology development, and formal analysis; **Muhammad Amsyar Izmi:** Literature review, data collection, simulation of results, and validation; **Rizauddin Saian:** Language editing and contribution to the interpretation of results.

9. REFERENCES

- Aulia, L., Wulandari, R., Diana, A. F., & Sulistijanti, W. (2024). Application of fuzzy AHP and new fuzzy Servqual methods for performance evaluation based on perceptions of service quality and satisfaction towards the implementation of educational programs. *Advances in Social Science, Education and Humanities Research*, 856, 584-589. https://doi.org/10.2991/978-2-38476-273-6_63
- Aziz, K. A. A., Idris, M. F. I. M., Daud, W. S. W., & Aziz, N. S. F. (2023). Nasyid competition assessment using fuzzy evaluation method. *Journal of Computing Research and Innovation*, 8(2), 12-19. <https://doi.org/10.24191/jcrinn.v8i2.351>
- Aziz, K. A. A., Idris, M. F. I. M., Daud, W. S. W., & Fauzi, M. M. (2023). Application of fuzzy analytic hierarchy process (FAHP) for the selection of best student award. *Journal of Computing Research and Innovation*, 8 (2), 80-90. <https://doi.org/10.24191/jcrinn.v8i2.345>
- Baker, J., Côté, J., & Abernethy, B. (2003). Sport-specific practice and the development of expert decision-making in team ball sports. *Journal of Applied Sport Psychology*, 15(1), 12–25. <https://doi.org/10.1080/10413200305400>
- Bozanic, D., & Pamucar, D. (2013). Modification of the analytical hierarchical process method and its application in decision making in the defense system. *Technology*, 68(2), 327-334.
- Bulent, O. M., Ugur, O., & Ozkan, B. (2017). Evaluation of sport mental toughness and psychological wellbeing in undergraduate student athletes. *Educational Research and Reviews*, 12(8), 483–487. <https://doi.org/10.5897/ERR2017.3216>
- Covassin, T., Beidler, E., Ostrowski, J., & Wallace, J. (2015). Psychosocial aspects of rehabilitation in sports. *Clinics in Sport Medicine*, 34(2), 199–212. <https://doi.org/10.1016/j.csm.2014.12.004>
- Denison, J., Mills, J. P., & Konoval, T. (2017). Sports' disciplinary legacy and the challenge of coaching

- differently. *Sport, Education and Society*, 22(6), 772–783. <https://doi.org/10.1080/13573322.2015.1061986>
- Diejomaoh, S. O. E., Akarah, E., & Tayire, F. O. (2015). Availability of facilities and equipment for sports administration at the local government areas of Delta State, Nigeria. *Academic Journal of Interdisciplinary Studies*, 4(2), 307–312. <https://doi.org/10.5901/ajis.2015.v4n2p307>
- Graham, d. n., & Burns, g. n. (2020). athletic identity and moral development: An examination of collegiate athletes and their moral foundations. *International Journal of Sport Psychology*, 51(2), 122–140. <https://doi.org/10.7352/IJSP.2020.51.122>
- Ibrahim, H. I., Jaafar, A. H., Muhammad, & Isa, A. (2016). Motivational climate, self-confidence, and perceived success among student athletes. *Procedia Economics and Finance*, 35, 503–508. [https://doi.org/10.1016/S2212-5671\(16\)00062-9](https://doi.org/10.1016/S2212-5671(16)00062-9)
- Idris, M. F. I. M., Aziz, K. A. A., & Aziz, N. S. F. A. (2023). selecting the effective ways to prevent Covid-19 from spreading using fuzzy AHP method. *Journal of Computing Research and Innovation*, 8(2), 112–123. <https://doi.org/10.24191/jcrinn.v8i2.348>
- Kubiak, C. (2012). *Perceived factors influencing athletic performance across career stages*. Digitala Vetenskapliga Arkivet. <https://www.diva-portal.org/smash/record.jsf?pid=diva2%3A619224&dsid=-7768>
- Macquet, A.-C., & Skalej, V. (2015). Time management in elite sports: How do elite athletes manage time under fatigue and stress conditions?. *Journal of Occupational and Organizational Psychology*, 88(2), 341–363. <https://doi.org/10.1111/joop.12105>
- Nacar, E., Gacar, A., Karahüseyinoğlu, M., & Gündoğdu, C. (2013). Analysis for sports facilities in sports high schools in terms of quality and quantity: Central Anatolia region sample. *Australian Journal of Basic and Applied Sciences*, 7(2), 627–631. <http://www.ajbasweb.com/old/ajbas/2013/February/627-631.pdf>
- Peng G., Han L., Liu Z., Guo Y., Yan J. & Jia X. (2021). An application of fuzzy analytic hierarchy process in risk evaluation model. *Frontier in Psychogy*. 12, 715003. <https://doi.org/10.3389/fpsyg.2021.715003>
- Rashid, N. S. A., Amin, N. A. M. & Mahad, N. F. (2020). Application of fuzzy analytic hierarchy process for contractor selection problem. *Gading Journal of Science and Technology*, 3(2), 109–117.
- Saaty, T. L. (1980). *The analytic hierarchy process*, New York: McGraw-Hill. <https://doi.org/10.21236/ADA214804>
- Saaty, T. L. & Ozdemir, M. S. (2015). How many judges should there be in a group?. *Annals of Data Science*, 1(3), 359–368. Springer. <https://10.1007/s40745-014-0026-4>.
- Sresto, M. A., Siddika, S. Haque, M. N. & Saroar, M. (2021). Application of fuzzy analytic hierarchy process and geospatial technology to identify groundwater potential zones in north-west region of Bangladesh. *Environmental Challenges*, 5, 100214. <https://doi.org/10.1016/j.envc.2021.100214>
- Suyudi, I. (2023). The digital revolution in sports: Analyzing the impact of information technology on athlete training and management. *Golden Ratio of Mapping Idea and Literature Format*, 3(2), 140–155. <https://doi.org/10.52970/grmilf.v3i2.343>
- Van Laarhoven, P.J.M., & Pedrycz, W. (1983). A fuzzy extension of Saaty's priority theory. *Fuzzy Sets and* <https://doi.org/10.24191/jcrinn.v10i2.520>

Systems, 11(1-3), 229-241. [https://doi.org/10.1016/S0165-0114\(83\)80082-7](https://doi.org/10.1016/S0165-0114(83)80082-7)

Wilson, R. S., David, G. K., Murphy, S. C., Angilletta, M. J. Jr., Niehaus, A. C., Hunter, A. H., & Smith, M. D. (2017). Skill, not athleticism, predicts individual variation in match performance of soccer players. *Proceedings of the Royal Society B: Biological Sciences*, 284(1868), 20170953. <https://doi.org/10.1098/rspb.2017.0953>

Zabihi, H., Alizadeh, M., Wolf, I. D., Karami, M., & Ahmad, A. (2020). A GIS-based fuzzy-analytic hierarchy process (F-AHP) for ecotourism suitability decision making : A case study of Babol in Iran. *Tourism Management Perspectives*, 36, 100726. <https://doi.org/10.1016/j.tmp.2020.100726>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

A Fuzzy Analytic Hierarchy Process (FAHP) Approaches for Investment Decision-Making in Malaysia

Mohd Fazril Izhar Mohd Idris^{1*}, Khairu Azlan Abd Aziz², Ardini Athirah Mhd Munawar³

^{1,2,3}*Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Perlis Branch, Arau Campus, 02600 Arau Perlis, Malaysia.*

ARTICLE INFO

Article history:

Received 23 May 2025

Revised 21 August 2025

Accepted 26 August 2025

Published 1 September 2025

Keywords:

Fuzzy Analytic Hierarchy Process

Investment Decision-Making

Multi-Criteria Decision-Making

Financial Literacy

Risk Assessment

Gold

DOI:

10.24191/jcrinn.v10i2.533

ABSTRACT

Investment decisions are essential for achieving financial growth and stability. Young Malaysians, however, often face challenges such as limited financial resources, rising living costs, and low financial literacy. Common investment options in Malaysia include gold, stocks, property, and cryptocurrency. Each option differs in capital requirements, profitability, risk level, and long-term viability. Selecting the most appropriate investment is difficult because decision making is often subjective and influenced by uncertainty. Conventional multi criteria decision making (MCDM) methods such as Simple Additive Weighting (SAW), Technique for Order of Preference by Similarity to Ideal Solution (TOPSIS), and Analytic Hierarchy Process (AHP) have been widely used, but they are less effective in dealing with this challenge. To overcome this limitation, this study applies the Fuzzy Analytic Hierarchy Process (FAHP), which integrates fuzzy logic with AHP to capture expert evaluations more realistically. Four investment alternatives are assessed based on capital, profit, risk, and sustainability. The results indicate that gold is the most preferred investment option, followed by property, stocks, and cryptocurrency. By applying FAHP to investment decision-making in Malaysia, this study introduces a novel framework that not only supports young investors in making strategic choices but also offers insights that may inform financial literacy programs and policy initiatives under uncertain market conditions.

1. INTRODUCTION

Investment is a fundamental component of personal financial management. It enables individuals to grow wealth, reduce the effects of inflation, and achieve long-term financial security. In Malaysia, popular investment options include gold, stocks, property, and cryptocurrency. Each alternative presents different levels of risk, return potential, capital requirement, and sustainability (Mushaddik & Nori, 2023; Hassan et al., 2022; Eaw et al., 2024). Despite this availability, many Malaysians — especially young adults and fresh graduates — find it difficult to make informed decisions. Limited financial resources, rising living costs,

^{1*} Corresponding author. *E-mail address:* fazrilizhar@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.533>

and insufficient financial literacy continue to restrict their ability to participate in suitable investment opportunities (Idris et al., 2013; Anuwar et al., 2024).

The urgency of enhancing investment literacy is underscored by concerning statistics. According to the Employees Provident Fund (EPF), as of 2021, only 3% of Malaysians are financially prepared for retirement, highlighting a critical gap in personal financial planning and investment practices (Malay Mail, 2021). This issue is exacerbated by the tendency of individuals to engage in high-risk or unsuitable investment schemes due to a lack of understanding regarding essential factors such as capital, profit, risk, and sustainability (Raut, 2020). Consequently, there is a pressing need for structured decision-making frameworks that can guide individuals, especially the younger generation, in selecting suitable investment options aligned with their financial goals and risk tolerance.

To address this gap, numerous multi-criteria decision-making (MCDM) methods have been utilized in previous studies to evaluate and rank investment alternatives. Techniques such as Simple Additive Weighting (SAW), Technique for Order of Preference by Similarity to Ideal Solution (TOPSIS), and the Analytic Hierarchy Process (AHP) have been widely applied in investment decision-making (Aliyeva, 2024; Dweiri & Al-Oqla, 2006; Purwita & Subriadi, 2019; Sahore, 2017). While these methods offer structured evaluation mechanisms, they often fall short in handling the vagueness and subjectivity inherent in human judgments, especially when dealing with qualitative criteria (Hodosi et al., 2023).

To overcome these limitations, the Fuzzy Analytic Hierarchy Process (FAHP) integrates fuzzy set theory with AHP, allowing decision-makers to express preferences using linguistic variables and accommodate the inherent ambiguity of human assessments (Sachdeva et al., 2023; Singh, 2016). FAHP has been successfully applied in various domains where subjective expert opinions are crucial, offering a more realistic and flexible decision-making framework compared to traditional methods.

Given the limited application of FAHP in the context of investment decision-making in Malaysia, this study aims to bridge this gap by applying FAHP to evaluate four major investment options, namely gold, stocks, property, and cryptocurrency, based on the criteria of capital, profit, risk, and sustainability. The difficulty of investment decision-making in Malaysia is compounded by several factors, including limited financial literacy among young adults, rising living costs, and the tendency to rely on subjective or uncertain judgments when assessing alternatives. Conventional MCDM approaches, while widely used, often fail to adequately address this ambiguity. By incorporating fuzzy logic into AHP, the FAHP method allows expert judgments to be expressed more realistically and reduces the effects of vagueness in evaluation. In this way, the study not only fills a methodological gap but also provides practical insights that can help young Malaysian investors make informed, strategic, and resilient investment decisions aligned with their financial capacity and long-term goals.

2. METHODOLOGY

2.1 Overview of FAHP applications

The Fuzzy Analytic Hierarchy Process (FAHP) has been widely applied across various fields to support complex decision-making, particularly when subjective judgments and qualitative factors are involved. Recent studies have demonstrated the effectiveness of FAHP in addressing problems that require multi-criteria evaluation. For instance, Mohd Idris et al. (2019) applied FAHP to determine the underlying causes of divorce in Perlis, showcasing the method's ability to prioritize qualitative social factors. Similarly, Abd Aziz et al. (2024a) employed FAHP to analyze factors influencing career choices among graduate students, providing insights into personal and economic considerations in academic decision-making. In the public health domain, Idris et al. (2023) utilized FAHP to select the most effective measures for preventing the spread of COVID-19, emphasizing its applicability in crisis management scenarios.

Furthermore, Idris et al. (2024) applied FAHP to evaluate courier service preferences among Malaysian users, while Zahrin et al. (2022) used the method to identify criteria for choosing tour packages in Langkawi Island. In the context of consumer behavior, Abd Aziz et al. (2024b) examined factors affecting online shopping platform selection, and Abd Aziz et al. (2023) utilized FAHP to select recipients of the Best Student Award. These diverse applications highlight FAHP's versatility and robustness in handling decision-making problems that involve ambiguity and subjective human judgments, thereby justifying its suitability for investment decision-making in this study.

2.2 Research framework and FAHP process

The research framework in this study serves as a systematic guideline for evaluating investment alternatives using the Fuzzy Analytic Hierarchy Process (FAHP). By integrating multi-criteria decision-making (MCDM) principles with fuzzy logic, the framework effectively addresses the ambiguity and subjectivity inherent in human judgments. This is particularly important when dealing with qualitative factors that influence investment decisions, where traditional crisp values may not capture expert preferences accurately.

The primary objective of this study is to assess and rank four common investment options in Malaysia, which are gold, stocks, property, and cryptocurrency. These alternatives are evaluated based on four key criteria, namely capital, profit, risk, and sustainability. The relationships between the overall goal, evaluation criteria, and investment alternatives are illustrated in the following Fig. 1.

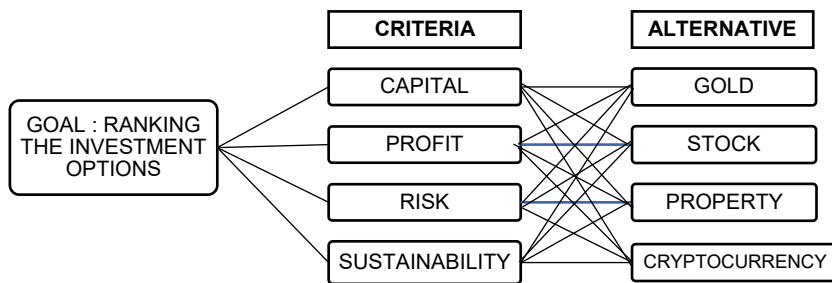


Fig. 1. Investment decision hierarchy tree

To operationalize this evaluation, the research framework adopts a structured sequence of steps guided by the FAHP methodology. This step-by-step process ensures that expert judgments are systematically collected, processed, and analyzed, while also accommodating the inherent uncertainty of subjective assessments. The general flow of the FAHP implementation in this study is illustrated in the following Fig. 2.

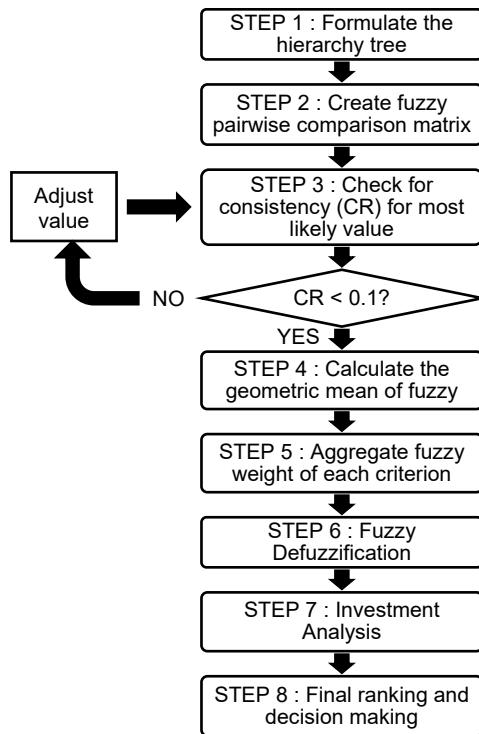


Fig. 2. FAHP process flowchart

The process begins with formulating the hierarchy tree, where the decision problem is structured into three levels: the overall goal, evaluation criteria, and available investment alternatives, as illustrated in Fig. 1. Following this, a fuzzy pairwise comparison matrix is constructed by gathering expert judgments through linguistic variables, which are represented as triangular fuzzy numbers (TFNs) to accommodate the inherent uncertainty and subjectivity in human assessments.

Triangular Fuzzy Numbers (TFNs) are a type of fuzzy number commonly used in decision-making to represent uncertain or imprecise judgments. A TFN is defined by three values: the lower (l), middle (m), and upper (u) bounds, written as (l, m, u) . These values form a triangular shape when plotted, representing the range and most likely value of an expert's opinion. TFNs are especially useful for translating linguistic terms (like "somewhat important" or "strongly important") into numerical values, allowing subjective judgments to be processed mathematically in fuzzy logic models like FAHP.

The membership function of a TFN, μ_{TFN} describes the degree to which each value within the range belongs to the fuzzy set, and is mathematically defined as follows:

$$\mu_{TFN} = \begin{cases} \frac{x-l}{m-l} & ; \quad l \leq x \leq m \\ \frac{u-x}{u-m} & ; \quad m \leq x \leq u \\ 0 & ; \quad otherwise \end{cases} \quad (1)$$

At this stage, the decision-maker refers to Table 1 to guide the selection of appropriate linguistic terms when performing pairwise comparisons between criteria and among alternatives. These qualitative

judgments are then translated into numerical values based on the standard AHP scale (ranging from 1 to 9). The resulting inputs are used to construct pairwise comparison matrices (PCMs), which systematically capture the expert's preferences. The general structure of the PCM, A^k based on the 1-9 scale, is represented in Eq. (2).

$$A^k = \begin{bmatrix} 1 & d_{12}^k & \cdots & d_{1n}^k \\ d_{21}^k & 1 & & d_{2n}^k \\ \vdots & & \ddots & \vdots \\ d_{n1}^k & d_{n2}^k & \cdots & 1 \end{bmatrix} \quad (2)$$

Here, k represents the k -th expert, and $d_{11}^k, d_{12}^k, \dots, d_{nn}^k$ denote the expert's pairwise comparison inputs for each factor, using a numerical scale from 1 to 9. It is important to note that these matrices are reciprocal in nature, meaning $d_{ji}^k = 1/d_{ij}^k$, and $d_{ii}^k = 1$ for all $i, j = 1, 2, \dots, n$. In other words, if the expert's preference d_{ij}^k is recorded in the upper triangle of the matrix, the corresponding reciprocal value d_{ji}^k will appear in the lower triangle. This reciprocal structure ensures the matrix consistently and accurately reflects the expert's judgments on the relative importance of each factor.

Table 1. Linguistic variables, along with their associated crisp scales and corresponding TFNs, used for pairwise comparisons

Linguistic Variables	AHP Scale (Non-Fuzzy Numbers)	Fuzzy AHP Scale (TFNs)	Fuzzy AHP Scale (Reciprocal TFNs)
Equally Important	1	(1,1,1)	(1,1,1)
Somewhat Important	3	(2,3,4)	(1/4,1/3,1/2)
Fairly Important	5	(4,5,6)	(1/6,1/5,1/4)
Strongly Important	7	(6,7,8)	(1/8,1/7,1/6)
Very Strongly Important	9	(9,9,9)	(1/9,1/9,1/9)
Intermediate Value	2	(1,2,3)	(1/3,1/2,1)
	4	(3,4,5)	(1/5,1/4,1/3)
	6	(5,6,7)	(1/7,1/6,1/5)
	8	(7,8,9)	(1/9,1/8,1/7)

To ensure the reliability of the expert evaluations, the consistency ratio (CR) is then calculated, verifying the logical consistency of the pairwise comparisons. Checking the consistency of the judgment matrix is essential, given the complexity of the decision-making process and the subjective nature of expert evaluations. To assess the reliability of the pairwise comparisons, the CR is computed using the following formula:

$$CR = \frac{CI}{RI} \quad (3)$$

where

$$CI = \frac{\lambda_{max} - n}{n - 1} \quad (4)$$

In this calculation, λ_{max} refers to the maximum eigenvalue of the comparison matrix, while n represents the size of the matrix. The Random Index (RI), which reflects the average consistency of randomly generated matrices, is presented in Table 2.

Table 2. Random consistency index (RI) values for different matrix sizes

Matrix Order, n	1	2	3	4	5	6	7	8	9	10
Random Index, RI	0.00	0.00	0.58	0.90	1.12	1.24	1.32	1.41	1.45	1.49

If the CR value is below 0.1 (or 10%), the matrix is considered to be consistent, and the evaluations are deemed acceptable. However, if the CR exceeds 0.1, the assessor is required to review and revise the matrix to improve consistency. It is important to note that the consistency check is performed using the crisp (non-fuzzy) version of the pairwise comparison matrix, as illustrated in Eq. (2).

Once the consistency of the experts' judgments has been verified, the pairwise comparison values are converted into Triangular Fuzzy Numbers (TFNs) based on the equivalence scale provided in Table 1. This conversion results in the formation of the Fuzzy Pairwise Comparison Matrix (FPCM) as illustrated below:

$$\tilde{A}^k = \begin{bmatrix} (1,1,1) & \tilde{d}_{12}^k & \cdots & \tilde{d}_{1n}^k \\ \tilde{d}_{21}^k & (1,1,1) & \cdots & \tilde{d}_{2n}^k \\ \vdots & \vdots & \ddots & \vdots \\ \tilde{d}_{n1}^k & \tilde{d}_{n2}^k & \cdots & (1,1,1) \end{bmatrix} \quad (5)$$

Here, $\tilde{d}_{11}^k, \tilde{d}_{12}^k, \dots, \tilde{d}_{nn}^k$ denote the input values that have been converted from crisp (non-fuzzy) numbers into Triangular Fuzzy Numbers (TFNs), as shown in Table 1.

Next, the geometric mean of the fuzzy judgments is calculated to aggregate the assessments from multiple experts into a single set of unified fuzzy weights for each criterion. These fuzzy weights are then combined with the fuzzy evaluations of the investment alternatives. The result is an aggregated fuzzy weight that reflects the overall importance of each investment alternative with respect to the selected criteria.

The following shows the calculation to obtain the average value for the updated matrix elements.

$$\tilde{A} = [\tilde{d}_{ij}] = \frac{\sum_{k=1}^K \tilde{d}_{ij}^k}{K} \quad (6)$$

where K denotes the total number of experts. In this context, $\tilde{d}_{ij}^k = (l_{ij}^k, m_{ij}^k, u_{ij}^k)$ where l , m , and u represent the lower, middle, and upper values of the Triangular Fuzzy Numbers (TFNs), respectively.

The fuzzy geometric mean, $\tilde{r}_i$ for the fuzzy comparison values of each criterion i is determined using Eq. (7).

$$\tilde{r}_i = (\tilde{d}_{i1} \times \tilde{d}_{i2} \times \cdots \times \tilde{d}_{iN})^N \quad (7)$$

Here, N represents the total number of criteria or alternatives. Next, the vector summation of the geometric means, $\sum_{i=1}^N \tilde{r}_i$, is computed, followed by calculating its reciprocal, $(\sum_{i=1}^N \tilde{r}_i)^{-1}$, using Eq. (8) and Eq. (9), respectively.

$$\sum_{i=1}^N \tilde{r}_i = \left(\sum l_{\tilde{r}_i}, \sum m_{\tilde{r}_i}, \sum u_{\tilde{r}_i} \right) \quad (8)$$

$$\left(\sum_{i=1}^N \tilde{r}_i \right)^{-1} = \left(\frac{1}{\sum u_{\tilde{r}_i}}, \frac{1}{\sum m_{\tilde{r}_i}}, \frac{1}{\sum l_{\tilde{r}_i}} \right) \quad (9)$$

To calculate the fuzzy weight for each criterion $\tilde{w}_i$, each geometric mean $\tilde{r}_i$ is multiplied by the reciprocal of the total vector summation of geometric means, as outlined in Eq. (10).

$$\tilde{w}_i = \tilde{r}_i \times \left(\sum_{i=1}^N \tilde{r}_i \right)^{-1} \quad (10)$$

The defuzzification process is subsequently carried out to transform the fuzzy weights into crisp numerical values, allowing for a clear and objective comparison among the alternatives. The calculation method for this process is shown as follows.

$$M_i = \frac{\tilde{l}_i + \tilde{m}_i + \tilde{u}_i}{3} \quad (11)$$

Next, the normalized value of M_i must be calculated to obtain the normalized weight Z_i , which represents the final weight. This is achieved using the following method:

$$Z_i = \frac{M_i}{\sum_{i=1}^N M_i} \quad (12)$$

This process is not limited to obtaining the final weight values for the criteria alone but is also extended to determine the final weights for each alternative. The same computational steps, as outlined from Eq. (2) to Eq. (12), are applied. However, while the final weight for each criterion is calculated only once, the final weights for the alternatives must be computed separately for each criterion. In other words, the weighting process for alternatives is repeated based on the total number of criteria.

Once the final weights for all criteria and alternatives have been obtained, the final step in the FAHP method is to calculate the total score for each alternative, denoted as T_A^n , using the following formula:

$$T_A^n = \sum_{i=1}^N Z_{A_i}^n \times Z_{C_i} \quad (13)$$

Here, n represents the n -th alternative, N is the number of criteria, $Z_{A_i}^n$ is the final weight of the alternative with respect to a given criterion, and Z_{C_i} is the final weight of the corresponding criterion.

Based on these values, an investment analysis is performed to compute the overall scores for each investment option. Finally, the most suitable investment alternative is determined through a final ranking and decision-making process, providing a structured and reliable basis for investment selection.

By following these systematic steps, the research framework provides a comprehensive approach for evaluating investment alternatives, ensuring robust and reliable decision-making outcomes that reflect expert insights despite inherent uncertainties.

3. RESULTS AND DISCUSSION

Data were collected using structured questionnaires distributed to two financial advisors, each with over ten years of professional experience in providing investment recommendations, portfolio planning, and risk assessment for clients. In this study, an individual was considered an expert if they possessed substantial professional experience in investment advisory, practical knowledge of financial markets, and active involvement in guiding clients' decision-making. This criterion ensures that the pairwise comparisons are grounded in reliable, informed, and contextually relevant judgments. The experts were asked to provide pairwise comparisons of the evaluation criteria as well as the investment alternatives, using linguistic terms that were subsequently converted into Triangular Fuzzy Numbers (TFNs) for analysis.

Each expert was provided with a complete set of questionnaires that required them to perform pairwise comparisons among all evaluation criteria and among all investment alternatives with respect to the four criteria: capital, profit, risk, and sustainability. Consequently, each expert generated a total of five pairwise comparison matrices (PCMs).

To ensure the reliability of the data, each PCM was subjected to a consistency check by calculating the Consistency Ratio (CR). A matrix was considered reliable if its CR value was below 0.1, which is the threshold recommended in AHP literature. All matrices in this study met this requirement, confirming that the expert judgments were logically consistent.

The validity of the data was ensured through three measures. First, experts were selected based on substantial professional experience in investment advisory, ensuring the judgments reflected domain knowledge. Second, the use of linguistic variables converted into Triangular Fuzzy Numbers (TFNs) allowed subjective judgments to be expressed more realistically, reducing cognitive bias. Third, the application of defuzzification and normalization produced crisp, comparable values, thereby enhancing the interpretability and accuracy of the results. Together, these steps ensured that the collected data were both valid and reliable for subsequent analysis.

3.1 Consistency test

Table 3 through Table 7 present the input data provided by Expert 1 (Exp1) and Expert 2 (Exp2), along with the Consistency Ratio (CR) results for each PCM. All CR values were found to be less than 0.1, indicating an acceptable level of consistency in the judgments. If a CR value exceeds 0.1, the judgment would be considered inconsistent, necessitating a re-evaluation or replacement of the expert's input.

Table 3. Pairwise comparison matrix for criteria, including the consistency ratio for expert 1 and 2

Criteria	Capital		Profit		Risk		Sustainability	
	Exp1	Exp2	Exp1	Exp2	Exp1	Exp2	Exp1	Exp2
Capital	1	1	4	3	3	8	8	9
Profit	1/4	1/3	1	1	1/4	2	2	1
Risk	1/3	1/8	4	1/2	1	1	2	1/3
Sustainability	1/8	1/9	1/2	1	1/2	3	1	1
CR							0.076	0.069

Table 4. Pairwise comparison matrix for alternative with respect to capital, including the consistency ratio for expert 1 and 2

Alternative with respect to Capital	Gold		Stock		Property		Cryptocurrency	
	Exp1	Exp2	Exp1	Exp2	Exp1	Exp2	Exp1	Exp2
Gold	1	1	2	3	5	3	4	5
Stock	1/2	1/3	1	1	6	2	3	7
Property	1/5	1/3	1/6	1/2	1	1	2	4
Cryptocurrency	1/4	1/5	1/3	1/7	1/2	1/4	1	1
CR							0.087	0.079

Table 5. Pairwise comparison matrix for alternative with respect to profit, including the consistency ratio for expert 1 and 2

Alternative with respect to Profit	Gold		Stock		Property		Cryptocurrency	
	Exp1	Exp2	Exp1	Exp2	Exp1	Exp2	Exp1	Exp2
Gold	1	1	2	1	3	5	5	2
Stock	1/2	1	1	1	2	5	3	1/2
Property	1/3	1/5	1/2	1/5	1	1	6	8
Cryptocurrency	1/5	1/2	1/3	2	1/6	8	1	1
CR							0.087	0.079

Table 6. Pairwise comparison matrix for alternative with respect to risk, including the consistency ratio for expert 1 and 2

Alternative with respect to Risk	Gold		Stock		Property		Cryptocurrency	
	Exp1	Exp2	Exp1	Exp2	Exp1	Exp2	Exp1	Exp2
Gold	1	1	2	2	3	5	5	4
Stock	1/2	1/2	1	1	2	6	7	3
Property	1/3	1/5	1/2	1/6	1	1	6	2
Cryptocurrency	1/5	1/4	1/7	1/3	1/6	1/2	1	1
CR							0.072	0.087

Table 7. Pairwise comparison matrix for alternative with respect to sustainability, including the consistency ratio for expert 1 and 2

Alternative with respect to Sustainability	Gold		Stock		Property		Cryptocurrency	
	Exp1	Exp2	Exp1	Exp2	Exp1	Exp2	Exp1	Exp2
Gold	1	1	1	2	5	3	2	5
Stock	1	1/2	1	1	5	2	1/2	3
Property	1/5	1/3	1/5	1/2	1	1	1/8	6
Cryptocurrency	1/2	1/5	2	1/3	8	1/6	1	1
CR							0.079	0.087

3.2 Fuzzy geometric mean

Since all CR values were within acceptable limits, the next step was to convert the crisp values in the PCMs into Triangular Fuzzy Numbers (TFNs) using the scale provided in Table 1. Each expert's TFNs were then averaged to generate a combined input for further computation. The averaged TFNs for all PCMs are presented in Tables 8 to 12.

Table 8. Average of pairwise comparison matrix for criteria

Criteria	Capital	Profit	Risk	Sustainability
Capital	(1,1,1)	(2.5,3.5,4,5)	(4.5,5.5,6.5)	(8,8.5,9)
Profit	(0.225,0.292,0.417)	(1,1,1)	(0.6,1.125,1.667)	(1,1.5,2)
Risk	(0.181,0.229,0.321)	(1.667,2.25,3)	(1,1,1)	(0.625,1.167,1.75)
Sustainability	(0.111,0.118,0.127)	(0.667,0.75,1)	(1.167,1.75,2.5)	(1,1,1)

Table 9. Average of pairwise comparison matrix for alternative with respect to capital

Alternative w.r.t capital	Gold	Stock	Property	Cryptocurrency
Gold	(1,1,1)	(1.5,2.5,3.5)	(3,4,5)	(3.5,4.5,5.5)
Stock	(0.292,0.417,0.75)	(1,1,1)	(3,4,5)	(4,5,6)
Property	(0.208,0.267,0.375)	(0.238,0.333,0.6)	(1,1,1)	(2,3,4)
Cryptocurrency	(0.183,0.225,0.292)	(0.188,0.238,0.333)	(0.267,0.375,0.667)	(1,1,1)

Table 10. Average of pairwise comparison matrix for alternative with respect to profit

Alternative w.r.t profit	Gold	Stock	Property	Cryptocurrency
Gold	(1,1,1)	(1,1.5,2)	(3,4,5)	(2.5,3.5,4.5)
Stock	(0.667,0.75,1)	(1,1,1)	(2.5,3.5,4.5)	(1.167,1.75,2.5)
Property	(0.208,0.267,0.375)	(0.25,0.35,0.625)	(1,1,1)	(2.556,3.063,3.571)
Cryptocurrency	(0.25,0.35,0.625)	(0.625,1.167,1.75)	(3.571,4.083,4.6)	(1,1,1)

Table 11. Average of pairwise comparison matrix for alternative with respect to risk

Alternative w.r.t risk	Gold	Stock	Property	Cryptocurrency
Gold	(1,1,1)	(1,2,3)	(3,4,5)	(3.5,4.5,5.5)
Stock	(0.333,0.5,1)	(1,1,1)	(3,4,5)	(4,5,6)
Property	(0.208,0.267,0.375)	(0.238,0.333,0.6)	(1,1,1)	(3,4,5)
Cryptocurrency	(0.183,0.225,0.292)	(0.188,0.238,0.333)	(0.238,0.333,0.6)	(1,1,1)

Table 12. Average of pairwise comparison matrix for alternative with respect to sustainability

Alternative w.r.t sustainable	Gold	Stock	Property	Cryptocurrency
Gold	(1,1,1)	(1,1.5,2)	(3,4,5)	(2.5,3.5,4.5)
Stock	(0.667,0.75,1)	(1,1,1)	(2.5,3.5,4.5)	(1.167,1.75,2.5)
Property	(0.208,0.267,0.375)	(0.25,0.35,0.625)	(1,1,1)	(2.556,3.063,3.571)
Cryptocurrency	(0.25,0.35,0.625)	(0.625,1.167,1.75)	(3.571,4.083,4.6)	(1,1,1)

Subsequently, the Fuzzy Geometric Mean was computed using Eq. (7), Eq. (8), and Eq. (9). The results, including the vector summation of the geometric mean and their reciprocals, are displayed in Tables 13 and 14. Table 13 shows the vector summation of the geometric mean for the criteria, while Table 14 displays the corresponding results for all alternatives with respect to each criterion.

Table 13. Vector summation of geometric mean of criteria

Criteria	TFN
Vector summation	(4.887,5.922,7.004)
Reciprocal	(0.143,0.169,0.205)

Table 14. Vector summation of geometric mean of alternatives with respect to all criteria

Alternative w.r.t.	Capital	Profit	Risk	Sustainability
Vector summation	(4.231,5.384,6.789)	(4.304,5.472,6.876)	(4.136,5.366,6.875)	(4.304,5.472,6.876)
Reciprocal	(0.147,0.186,0.236)	(0.145,0.183,0.232)	(0.145,0.186,0.242)	(0.145,0.183,0.232)

The process continued with the calculation of fuzzy weights for all criteria and alternatives, as described in Eq. (10). These fuzzy weights were then defuzzified and normalized using Eq. (11) and Eq. (12). The results are presented in Table 15 (for criteria) and Table 16 (for alternatives with respect to all criteria).

Table 15. Fuzzy weight, defuzzified and normalized value for all criteria

Criteria	Fuzzy Weight	Defuzzified	Normalized	Rank
Capital	(0.44,0.604,0.826)	0.6233	0.5966	1
Profit	(0.087,0.142,0.223)	0.1502	0.1438	3
Risk	(0.094,0.149,0.234)	0.1589	0.1520	2
Sustainability	(0.077,0.106,0.154)	0.1125	0.1076	4
Total		1.0449	1.000	

Table 16. Fuzzy weight, defuzzified and normalized value for all alternative with respect to criteria

Criteria	Alternative	Fuzzy Weight	Defuzzified	Normalized
Capital	Gold	(0.293,0.482,0.739)	0.5046	0.4693
	Stock	(0.201,0.316,0.514)	0.3437	0.3197
	Property	(0.082,0.134,0.23)	0.1487	0.1383
	Cryptocurrency	(0.045,0.07,0.119)	0.0782	0.0727
Total			1.0752	1.0000
Profit	Gold	(0.24,0.392,0.601)	0.4109	0.3827
	Stock	(0.171,0.268,0.425)	0.2880	0.2683
	Property	(0.088,0.134,0.222)	0.1478	0.1376
	Cryptocurrency	(0.125,0.208,0.347)	0.2269	0.2114
Total			1.0736	1.0000
Risk	Gold	(0.261,0.456,0.729)	0.4820	0.4433
	Stock	(0.205,0.331,0.566)	0.3674	0.3379
	Property	(0.09,0.144,0.249)	0.1610	0.1481
	Cryptocurrency	(0.044,0.068,0.119)	0.0768	0.0707
Total			1.0872	1.0000
Sustainability	Gold	(0.24,0.392,0.601)	0.4109	0.3827
	Stock	(0.171,0.268,0.425)	0.2880	0.2683
	Property	(0.088,0.134,0.222)	0.1478	0.1376
	Cryptocurrency	(0.125,0.208,0.347)	0.2269	0.2114
Total			1.0736	1.0000

3.3 Final ranking of investment alternatives

The final step was to determine the ranking of each investment alternative. As shown in Table 17, the total score for each alternative was computed by multiplying its normalized value with the normalized weight of each criterion, then summing the products. The alternative with the highest total score was ranked first, followed by the next highest, and so on. This ranking process is illustrated in Eq. (13).

Table 17. Ranking of the alternative

Alternative	Capital	Profit	Risk	Sustainability	Total (Summation the product of alternative and criteria)	Rank
	0.5966	0.1438	0.1520	0.1076		
Gold	0.4693	0.3827	0.4433	0.3827	0.4436	1
Stock	0.3197	0.2683	0.3379	0.2683	0.3095	2
Property	0.1383	0.1376	0.1481	0.1376	0.2359	3
Cryptocurrency	0.0727	0.2114	0.0707	0.2114	0.1605	4

Table 17 summarizes the final ranking of investment alternatives which include Gold, Stock, Property, and Cryptocurrency. The ranking is based on the total score calculated from the weighted sum of normalized fuzzy values across four evaluation criteria: capital, profit, risk, and sustainability. The criteria weights, derived through FAHP, are as follows: Capital (0.5966), Profit (0.1438), Risk (0.1520), and Sustainability (0.1076). Notably, Capital carries the highest weight, indicating that the experts considered it the most influential factor in investment decision-making.

Gold emerged as the top-ranked investment alternative, with the highest total score of 0.4436. This result reflects Gold's strong performance across all four criteria, especially under the heavily weighted capital category, where it scored 0.4693. Its relatively high values in profit, risk, and sustainability further contributed to its dominance. This suggests that Gold is perceived as a relatively secure and stable investment, making it favorable under current financial conditions in Malaysia.

Stock was ranked second, with a total score of 0.3095. It showed moderately strong values in capital (0.3197) and risk (0.3379), though it was slightly less competitive in profit (0.2683) and sustainability (0.2683) compared to Gold. Nonetheless, its balanced performance across all criteria indicates that stocks remain a viable investment option, albeit with slightly higher perceived risk and volatility than gold.

Property ranked third, with a total score of 0.2359. It consistently scored lower across all criteria, particularly in capital (0.1383) and profit (0.1376). Despite being traditionally viewed as a stable long-term investment, these findings suggest that the high capital requirements and slower liquidity may reduce its appeal, especially in comparison to more liquid assets like gold and stocks.

Cryptocurrency was ranked last, with the lowest total score of 0.1605. While it performed reasonably well in profit (0.2114) and sustainability (0.2114), its extremely low scores in capital (0.0727) and risk (0.0707) significantly affected its overall ranking. This outcome reflects the experts' cautious stance on cryptocurrencies, likely due to their high volatility, regulatory uncertainty, and speculative nature, making them less suitable for risk-averse investors in the Malaysian context.

The findings indicate that among the four alternatives, Gold is the most preferred investment, followed by Stock, Property, and Cryptocurrency. The strong emphasis placed on capital as the most critical factor influenced the rankings significantly. These insights can aid individual investors and financial planners in prioritizing investment options that align with their financial goals and risk tolerance.

4. CONCLUSIONS

This study employed the Fuzzy Analytic Hierarchy Process (FAHP) to support investment decision-making in the Malaysian context by evaluating and ranking four popular investment options namely gold, stocks,

property, and cryptocurrency. These options were assessed based on four essential criteria which are capital requirement, profit potential, risk level, and sustainability.

Through structured expert input and the integration of fuzzy logic, the methodology successfully addressed the vagueness and imprecision commonly found in human judgment, particularly when evaluating qualitative aspects. The analysis revealed that gold emerged as the most preferred investment alternative, followed by stocks, property, and cryptocurrency. This ranking was significantly influenced by the capital criterion, which had the greatest weight in the final decision-making process. Gold's high score reflects its historical reputation for financial stability and relatively lower risk, making it a favorable choice for conservative investors.

Overall, the FAHP method demonstrated its strength in handling complex, multi-criteria decisions by offering a structured and transparent evaluation process. The study contributes meaningful insights for individual investors and financial advisors seeking rational, data-driven approaches to selecting suitable investment avenues.

Based on the findings of this study, several recommendations can be drawn to support more informed investment decisions. Individual or young investors, particularly those with limited experience, are encouraged to consider gold and stocks as preferred investment options, as they demonstrated the highest rankings based on the selected evaluation criteria. However, investment choices should still be guided by personal financial goals, risk tolerance, and long-term objectives. Financial advisors may find the FAHP approach useful as a structured decision-support tool, enabling them to provide clients with more data-driven and transparent recommendations.

Future research is recommended to expand the number of evaluation criteria, such as liquidity, market volatility, and taxation, to reflect a more comprehensive investment landscape. Additionally, incorporating inputs from a larger and more diverse group of experts could enhance the accuracy and generalizability of the results. From a policy perspective, the application of FAHP could be extended to national financial literacy initiatives, offering tools and frameworks to guide the public in making sound, rational investment decisions.

5. ACKNOWLEDGEMENTS/FUNDING

The authors gratefully acknowledge the support of Universiti Teknologi MARA (UiTM), Perlis Branch, Arau Campus, Malaysia, for providing the necessary facilities and financial assistance for this research. The authors also extend their sincere appreciation to the anonymous reviewers for their valuable and constructive feedback, which greatly contributed to improving the quality of this study.

6. CONFLICT OF INTEREST STATEMENT

The authors declare that this research was conducted without any personal, commercial, or financial conflicts of interest, and that no conflicts exist with the funders.

7. AUTHOR'S CONTRIBUTIONS

Mohd Fazril Izhar Mohd Idris: Conceptualisation, supervision, methodology, formal analysis, investigation, writing- review and editing, and validation; **Khairu Azlan Abd Aziz:** Conceptualisation, methodology, and formal analysis; **Ardini Athirah Mhd Munawar:** Conceptualisation, methodology, formal analysis, and writing-original draft.

8. REFERENCES

- Abd Aziz, K. A., Daud, W. S. W., Idris, M. F. I. M., & Saimuddi, S. (2024a). Fuzzy analytical hierarchy process for analysing the factors that influence the career choices of graduate students at UiTM Perlis. *Journal of Computing Research and Innovation*, 9(1), 180-196. <https://doi.org/10.24191/jcrinn.v9i1>
- Abd Aziz, K. A., Daud, W. S. W., Idris, M. F. I. M., & Zakaria, S. N. (2024b). The ranking of factors in selecting the online shopping platform in Malaysia based on fuzzy analytic hierarchy process. *Journal of Computing Research and Innovation*, 9(1), 31-42. <https://doi.org/10.24191/jcrinn.v9i1>
- Abd Aziz, K. A., Idris, M. F. I. M., Daud, W. S. W., & Fauzi, M. M. (2023). Application of fuzzy analytic hierarchy process (FAHP) for the selection of best student award. *Journal of Computing Research and Innovation*, 8(2), 80-90. <https://doi.org/10.24191/jcrinn.v8i2.345>
- Aliyeva, K. R. (2024). Research on group decision making in capital investment based on fuzzy technique for order preference by similarity to ideal solution (TOPSIS). *Journal of Fuzzy Extension and Applications*, 5(3), 313-329. <https://doi.org/10.22105/jfea.2024.446437.1397>
- Anuwar, M. A. M., Adenan, F., Arif, M. I. A. M., Rosli, M. S. D. A., Ab Hamid, M. H., & Remly, M. R. (2024). Perception on financial literacy among youth: A study on UiTM Puncak Alam student. *Al-Qanatr: International Journal of Islamic Studies*, 33(3), 389-397. <https://al-qanatr.com/aq/article/view/881>
- Dweiri, F., & Al-Oqla, F. M. (2006). Material selection using analytic hierarchy process. *International Journal of Computer Applications in Technology*, 26(4), 182-189. <https://doi.org/10.1504/IJCAT.2006.010763>
- Eaw, H. C., Loebiantoro, I. Y., Jap, K. P., Shakur, E. S. A., & Voon, A. (2024, July). Green stock investment preferences among adult investors in east Malaysia. In *IOP Conference Series: Earth and Environmental Science* (Vol. 1372, No. 1, p. 012086). IOP Publishing. <https://doi.org/10.1088/1755-1315/1372/1/012086>
- Hassan, M. M., Ahmad, N., & Hashim, A. H. (2022). Housing property investment opportunity in Malaysia: Things that new investor should know. *International Journal of Academic Research in Business and Social Sciences*, 12(1), 924-941. <http://dx.doi.org/10.6007/IJARBSS/v12-i1/12003>
- Hodosi, G., Sule, E., & Bodis, T. (2023, November). Multi-criteria decision making: A comparative analysis. In *Economic and Social Development (Book of Proceedings)*, 103rd International Scientific Conference on Economic and Social Development (p. 81).
- Idris, F. H., Krishnan, K. S. D., & Azmi, N. (2013). Relationship between financial literacy and financial distress among youths in Malaysia-An empirical study. *Geografia-Malaysian Journal of Society and Space*, 9(4), 106-117. <http://ejournals.ukm.my/gmjss/article/download/18274/5760>
- Idris, M. F. I. M., Abd Aziz, K. A., & Abd Aziz, N. S. F. (2023). Selecting the effective ways to prevent covid-19 from spreading using fuzzy AHP method. *Journal of Computing Research and Innovation*, 8(2), 112-123. <https://doi.org/10.24191/jcrinn.v8i2.348>
- Idris, M. F. I. M., Abd Aziz, K. A., & Muhammad Munib, N. (2024). FAHP-based assessment of courier service preferences among Malaysian users. *Journal of Computing Research and Innovation*, 9(2), 376-395. <https://doi.org/10.24191/jcrinn.v9i2.459>
- Malay Mail. (2021, October 12). *EPF reveals only 3% of Malaysians can afford retirement*. MalayMail. <https://www.malaymail.com/>

- Mohd Idris, M. F. I., Abd Aziz, K. A., & Mohd Kamarulzaman, M. F. (2019). Determining the causes of divorce in Perlis using fuzzy analytic hierarchy process. *Jurnal Intelek*, 14(1), 56-65.
- Mushaddik, I. N., & Nori, A. W. J. (2023). Malaysia's potential revolution: Embracing gold-backed cryptocurrency into international net settlement via blockchain could transform economic and financial resilience. *Journal of Islam in Asia (E-ISSN 2289-8077)*, 20(3), 337-382. <https://doi.org/10.31436/jia.v20i3.1188>
- Purwita, A., & Subriadi, A. (2019). Literature review—using multi-criteria decision-making methods in information technology (IT) investment. In *Proceedings of the 1st International Conference on Business, Law And Pedagogy, ICBLP 2019*. European Alliance for Innovation (EAI). <https://doi.org/10.4108/eai.13-2-2019.2286076>
- Raut, R. K. (2020). Past behaviour, financial literacy and investment decision-making process of individual investors. *International Journal of Emerging Markets*, 15(6), 1243-1263. <https://doi.org/10.1108/IJOEM-07-2018-0379>
- Sachdeva, M., Lehal, R., Gupta, S., & Gupta, S. (2023). Influence of contextual factors on investment decision-making: a fuzzy-AHP approach. *Journal of Asia Business Studies*, 17(1), 108-128. <https://doi.org/10.1108/JABS-09-2021-0376>
- Sahore, A. (2017). Decision-making for investments in stocks using saw and topsis methods of multi-criteria decision-making. *NICE Journal of Business*, 12(2).
- Singh, B. (2016). Analytical hierarchical process (AHP) and fuzzy AHP applications-A review paper. *International Journal of Pharmacy and Technology*, 8(4), 4925-4946.
- Zahrin, K. A. D. K., Ab Halim, H. Z., Idris, M. F. I. M., Fauzi, N. F., Bakhtiar, N. S. A., & Khairudin, N. I. (2022). Determining the criteria of choosing tour packages in Langkawi Island using FAHP. *Journal of Computing Research and Innovation*, 7(2), 349-356. <https://doi.org/10.24191/jerinn.v7i2.330>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Evaluating User Acceptance of the UnityHub Web-based System in Higher Education: A Descriptive Analysis

Norazah Umar¹, Nurhafizah Ahmad^{2*}, Jamal Othman³, Rozita Kadar⁴

^{1,2,3,4}*Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA Pulau Pinang Branch, Permatang Pauh Campus, 13500 Permatang Pauh, Pulau Pinang, Malaysia.*

ARTICLE INFO

Article history:

Received 25 June 2025

Revised 28 August 2025

Accepted 28 August 2025

Published 1 September 2025

Keywords:

Web-Based System

Communication Platform

User Acceptance

Technology Acceptance Model

Information System

DOI:

10.24191/jcrinn.v10i2.549

ABSTRACT

The transition to hybrid learning in Malaysian higher education created a need for centralized platforms that enhance academic communication and class management. UnityHub was introduced as a web-based system to simplify group registration and strengthen student-lecturer interaction. The purpose of this study was to assess students' acceptance of UnityHub using the Technology Acceptance Model (TAM) and to examine both the overall level of acceptance and the relative strengths and weaknesses of its dimensions. A total of 90 students from Universiti Teknologi MARA, Penang Branch participated in the survey, which used a TAM-based questionnaire comprising 24 items. Data were analyzed using descriptive statistics, including mean scores and standard deviations. The findings show that satisfaction scored highest (mean = 4.16), followed by attitude (3.80) and ease of use (3.78). Perceived usefulness (3.71) and intention to use (3.72) were moderately high, while self-efficacy was lowest (3.42). Overall, UnityHub achieved strong acceptance for usability and efficiency, though improvements in academic integration and user confidence are needed. The study is limited to a single department sample; future research should involve multiple faculties and compare UnityHub with other platforms to enhance generalizability.

1. INTRODUCTION

Effective communication is essential in educational settings, particularly in online and hybrid learning environments where physical interactions are limited. The COVID-19 pandemic accelerated the transition to such learning models in Malaysian higher education, revealing critical challenges in sustaining engagement, managing academic information, and maintaining timely interactions between students and lecturers. In these contexts, the use of digital communication platforms becomes vital for supporting continuous learning and academic coordination.

However, fragmented communication channels and the lack of centralized systems often lead to delays, confusion, and disengagement. Recognizing these issues, UnityHub was developed as a web-based

^{2*} Corresponding author. *E-mail address:* nurha9129@uitm.edu.my
<https://doi.org/10.24191/jcrinn.v10i2.549>

academic management platform that organizes student-lecturer communication and class group registration into a single, accessible interface. Designed to streamline communication and information access, UnityHub enables students to join their respective class groups based on current timetables, receive timely updates from lecturers, and manage academic participation via mobile devices. While the platform addresses key communication and coordination gaps in hybrid learning, its long-term effectiveness depends heavily on user acceptance. Educational technologies regardless of their technical features must be perceived as useful and easy to use by students to ensure sustained adoption. This is particularly important in the post-pandemic context, where digital tools are expected to play a permanent role in academic processes. Therefore, understanding user perceptions is critical.

This study adopts the Technology Acceptance Model (TAM) to evaluate students' acceptance of UnityHub. TAM is a widely recognized theoretical framework used to predict users' acceptance of information systems based on two primary constructs: perceived usefulness (PU) and perceived ease of use (PEU). These variables influence users' attitudes, satisfaction, and intention to continue using a system. By incorporating TAM, this study provides a theoretically grounded analysis of UnityHub's reception among its target users. As digital platforms like UnityHub play a central role in academic communication, it is important to evaluate how students perceive and accept such systems. This study, guided by the Technology Acceptance Model (TAM), uses descriptive analysis to assess students' acceptance of UnityHub. Accordingly, the objectives of this research are as follows:

1. To assess students' acceptance of UnityHub across six key constructs self-efficacy, satisfaction, perceived usefulness, perceived ease of use, intention to use, and attitude toward use using descriptive statistical analysis.
2. To examine the relative strengths and weaknesses of UnityHub's acceptance dimensions in order to highlight areas of high student approval as well as aspects that indicate only moderate acceptance.

The subsequent sections explore relevant literature on the challenges of online communication, explain the development of the UnityHub web-based system, outline the research methodology, and present the analysis, results, and discussion. The final section concludes the research findings.

2. LITERATURE REVIEW

Technology integration in higher education has grown rapidly, particularly in response to the COVID-19 pandemic, which accelerated the adoption of online and hybrid learning platforms. This shift has brought user acceptance into sharper focus, as the success of any educational technology depends on how well users perceive and engage with it. Therefore, understanding the factors that influence users' acceptance and intention to use digital academic platforms is critical in evaluating system effectiveness beyond technical functionality.

The Technology Acceptance Model (TAM), first introduced by Davis (1989), remains one of the most established theoretical frameworks for evaluating information system adoption. TAM posits that two key constructs perceived usefulness (PU) and perceived ease of use (PEU) shape users' attitudes toward a system, which subsequently influence their behavioral intention to use it. A system is likely to be adopted when users believe it will improve their performance (PU) and is easy to operate (PEU). These constructions have been widely validated in studies of information systems and online applications across various domains.

In educational contexts, TAM has been applied extensively to examine students' interactions with e-learning platforms, digital portals, and web-based communication tools. For instance, Xie and Lee (2015) reported that both PU and PEU significantly influenced students' acceptance of social learning technologies.

<https://doi.org/10.24191/jcrinn.v10i2.549>

Hollister et al. (2022) observed that during pandemic-driven online learning, students evaluated system adoption largely through the lens of functionality and ease of navigation. Similarly, Bashir (2018) highlighted that user interface clarity and intuitive design are crucial for fostering positive learning experiences in multimedia courseware.

Beyond the original TAM constructs, researchers have incorporated additional variables to reflect the complex nature of user behavior. Among these, self-efficacy, the belief in one's ability to use a system effectively—has emerged as a consistent predictor of technology adoption (Meng & Zhang, 2023). Students who feel confident in their ability to navigate digital platforms without assistance are more likely to perceive the system as useful and user-friendly. Likewise, attitude toward use and intention to use are often included in extended TAM models to capture motivational and behavioral dimensions more comprehensively (Venkatesh et al., 2012). These expanded frameworks offer a more holistic understanding of how users form judgments about a system, particularly in educational environments.

Recent regional studies have confirmed the relevance of these constructs in the Malaysian higher education landscape. Al Husaini and Shukor (2023) identified time-efficiency, self-efficacy, and information access as key drivers of acceptance in student-facing systems. Abbas et al. (2019) found that social media platforms used for academic collaboration are adopted when they are perceived as both practical and engaging. These studies underscore the importance of designing systems that align with users' preferences, behaviors, and contextual needs.

In measuring acceptance, many studies rely on quantitative descriptive analysis of TAM-based surveys. Researchers typically use Likert-scale instruments to capture users' responses on PU, PEU, satisfaction, intention, and related constructs. Descriptive statistics such as means, standard deviations, and response frequencies are then used to interpret acceptance levels and highlight system strengths or weaknesses (Lee & Kozar, 2012; Rawashdeh et al., 2021). This approach provides both depth and breadth, offering a data-driven understanding of user reception and system usability.

Despite the extensive literature on TAM and technology acceptance, limited research has examined the adoption of communication-specific platforms like UnityHub, particularly within the Malaysian post-pandemic educational context. Moreover, few studies have integrated descriptive analysis with TAM constructs to evaluate systems developed for academic communication and group coordination. This study addresses that gap by applying TAM to investigate students' perceptions, satisfaction, and intention to use UnityHub, offering practical insights for improving digital communication infrastructure in higher education.

3. THE DEVELOPMENT OF UNITYHUB: A WEB-BASED SYSTEM FOR EFFECTIVE DISTANCE LEARNING MANAGEMENT

This section discusses the development of UnityHub, a web-based system designed to streamline the management of students in distance learning. Through this platform, lecturers can ensure students stay up to date with their coursework, minimizing the risk of them falling behind in online classes. UnityHub provides students with essential information from the start of the semester, simplifying the process by offering a list of registered subjects along with corresponding class groups based on the latest semester schedule. The user-friendly interface allows students to select their subjects and groups easily, enabling them to join designated groups with just a few clicks. Additionally, UnityHub is accessible via mobile devices, ensuring convenient access for all users.

The development of UnityHub involves three key stakeholders: academic management (including lecturers), students, and system developers. Fig. 1 illustrates the collaborative relationship among these parties. The process begins with the academic management team providing a timetable, which is then

distributed to lecturers and students. Based on this timetable, lecturers register their respective groups on platforms such as WhatsApp or Telegram to obtain the necessary platform link.

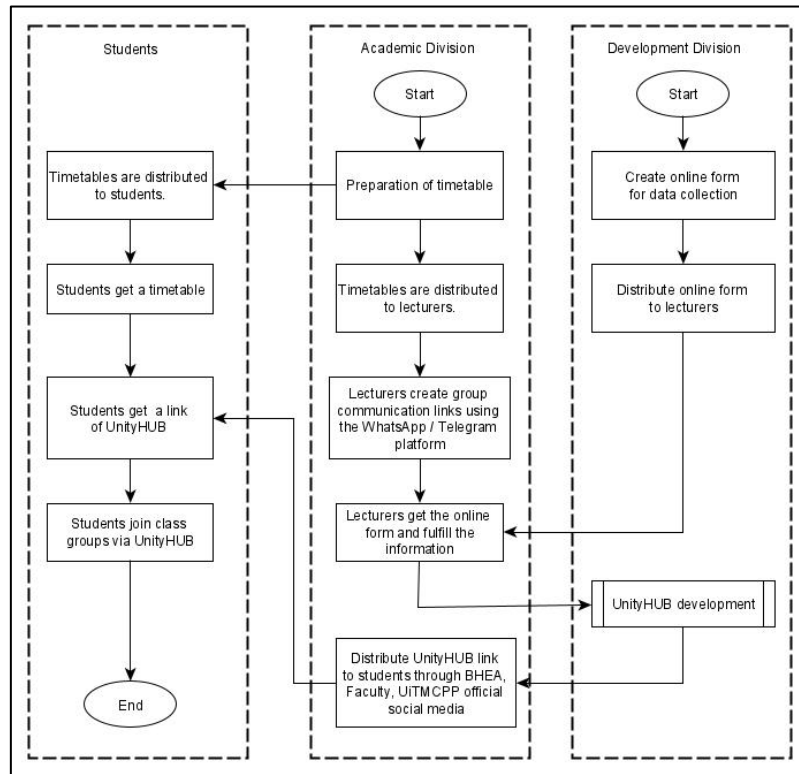


Fig. 1. Flowchart of class group registration using UnityHub

Simultaneously, the system development team creates an online form (Fig. 2) to collect essential data from lecturers for the UnityHub platform. This data includes information such as subject area, subject code, subject name, class group, lecturer's name, communication platform, and group link. Once the lecturers submit this information, the UnityHub development can proceed. When development is complete, the platform link is distributed to students through faculties and official social media channels, allowing students to join their respective groups based on the timetable simply by clicking the link.

Link Whatsapp Group Pelajar

Bidang:

☐ Matematik

☐ Statistik

☐ Sains Komputer

Kod Subjek (eg: CSC426) :

Your answer

Nama Subjek (eg: Introduction To Programming) :

Your answer

Kumpulan (eg: PHM1112A1):

Your answer

Nama Penderah:

Your answer

Platform:

☐ Whatsapp

☐ Telegram

Link Group:

Your answer

Submit

Fig. 2. Online form

UnityHub offers an efficient method for gathering geographically dispersed students, eliminating the need for lecturers to invite each student individually. Before the first class, lecturers often need to assemble their students promptly, and UnityHub streamlines this process, ensuring all students are organized into their respective groups within a week, which enables seamless classroom management.

Fig. 3 provides an overview of the UnityHub interface, which includes six tabs: HOME, SARJANA, SARJANA MUDA, DIPLOMA, PRA-DIPLOMA, and KENALI JSKM. The HOME section serves as a platform for the JSKM coordinator to communicate essential messages to users. Another tab displays a comprehensive list of subjects offered under JSKM, while the KENALI JSKM tab links users to the official JSKM website. For instance, Pre-diploma students (tag no.1) can click on the PRA-DIPLOMA tab to access a list of pre-diploma subjects. They can then select their subjects (tag no.2), and the platform will display associated class groups. By clicking the provided link, students can join their designated class group effortlessly. Both lecturers and students receive timely notifications via the registered platform (tag no.3), promoting effective communication and engagement.

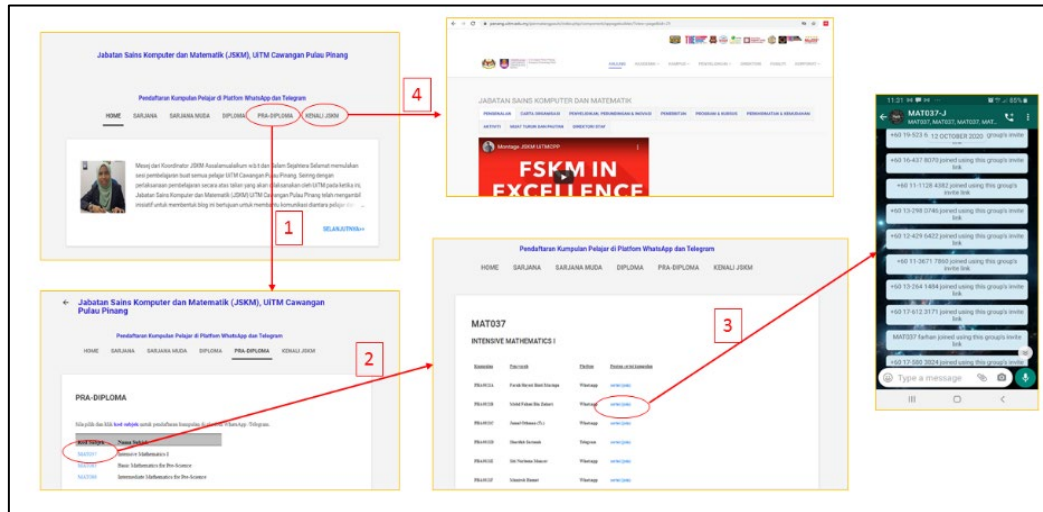


Fig. 3. UnityHub Interface

UnityHub is an innovative web-based system designed to enhance distance learning management. By simplifying the process of organizing students into class groups, UnityHub helps students stay on track with their coursework. Its user-friendly interface and mobile accessibility provide convenience for both lecturers and students. UnityHub effectively unites academic management, lecturers, and students, creating an efficient and streamlined distance learning experience. With its comprehensive features and seamless integration, UnityHub has the potential to revolutionize distance learning program management.

4. METHODOLOGY

4.1 Participants

The population of this study comprised undergraduate students from Universiti Teknologi MARA (UiTM), Penang Branch. A total of 90 students were selected using cluster sampling, with the sample drawn from five classes under lecturers in the Computer and Mathematics Department. These classes were selected based on their active engagement with the UnityHub system during the academic semester, thereby representing users with direct system exposure. While this sampling approach ensured data relevance, it introduces a potential bias, as the sample may not fully represent students from other faculties or campuses. This limitation is acknowledged and discussed further in the final section.

4.2 Instrument

Data was collected using a structured online questionnaire adapted from established instruments based on the Technology Acceptance Model (TAM) as proposed by Davis (1989) and later extended by Venkatesh et al. (2012). The questionnaire consisted of 24 items measuring six TAM-related constructs: self-efficacy, perceived usefulness (PU), perceived ease of use (PEU), satisfaction, intention to use, and attitude toward use. Items were rated on a five-point Likert scale ranging from 1 (“strongly disagree”) to 5 (“strongly agree”). An additional section captured demographic information, including gender, age, academic level, field of study, sources of UnityHub information, and reasons for using the system. To ensure content validity, the instrument underwent expert review by two senior lecturers with research expertise in educational technology and instrument design. Their feedback informed minor adjustments in item wording to ensure contextual relevance. A pilot test involving 15 students from a different faculty was

conducted to assess clarity and consistency of the items prior to full-scale distribution. Based on feedback, minor modifications were made to improve item phrasing without altering construct meanings.

4.3 Procedure

The data collection process was conducted through an online survey administered via a secure digital platform. An invitation to participate in the study was disseminated through official faculty communication channels, including email and WhatsApp group announcements coordinated by class lecturers. The message included a brief explanation of the study's purpose, a statement assuring confidentiality and voluntary participation, and a link to the online questionnaire. All selected participants were students enrolled in courses that actively utilized UnityHub during the semester. Prior to completing the survey, students were informed that their responses would be used solely for research purposes and would remain anonymous, with no identifying information being collected. Upon accessing the link, participants were directed to a brief introduction page outlining the scope of the study, followed by the main questionnaire. Each respondent was given an estimated time of 10 to 15 minutes to complete the survey at their own convenience.

4.4 Data Analysis

Data was analyzed using IBM SPSS Statistics Version 26. Descriptive statistics were used to address both research questions. For categorical data such as demographics and reasons for using UnityHub, frequencies and percentages were calculated. For Likert-scale items measuring TAM constructs, mean scores and standard deviations were computed to assess levels of user satisfaction, perceived usefulness, perceived ease of use, and behavioral intention. These descriptive metrics provided insight into the distribution of responses and the general acceptance level of UnityHub among students.

Instrument reliability was evaluated using Cronbach's Alpha for each construct. Following guidelines by Hinton et al. (2004), alpha values above 0.70 were interpreted as indicating high reliability, while scores above 0.90 signified excellent internal consistency. All constructs met or exceeded the 0.80 threshold, confirming the instrument's reliability for use in this study. To facilitate interpretation of the mean scores, student responses were categorized into five levels ranging from Very Low to Very High, as shown in Table 1. This classification provided a clearer understanding of whether each construct reflected low, moderate, or high levels of acceptance of UnityHub.

Table 1. Mean score interpretation

Mean Score	Interpretation
4.30 – 5.00	Very High
3.50 – 4.29	High
2.70 – 3.49	Moderate
1.90 – 2.69	Low
1.00 – 1.89	Very Low

The cut-off ranges in Table 1 follow the approach suggested by Zaki and Ahmad (2017), ensuring consistency with prior studies in interpreting students' responses. Based on this interpretation scale, the following section presents the descriptive results of the study. The findings are organized according to the TAM constructs, with each mean score reported together with its corresponding interpretation level to provide a clearer picture of students' acceptance of UnityHub.

5. RESULTS AND DISCUSSION

5.1 Descriptive analysis

The survey included 90 students, with 60% identifying as female and 40% as male, as illustrated in Fig. 4. Table 2 presents the demographic characteristics of the respondents. Most respondents (90%) were young adults aged between 18 and 22 years. In terms of academic qualifications, the majority were diploma students (63.3%), followed by degree students (35.6%), and only one respondent was a master's student. Most respondents were from fields related to Science and Technology (84.4%), while 15.6% were from Social Science.

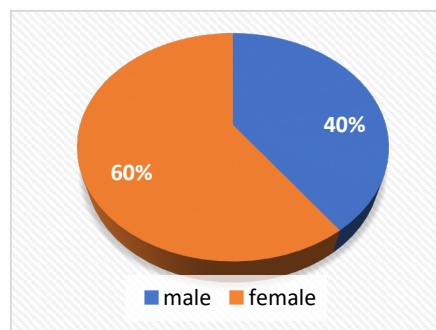


Fig. 4. Respondent gender

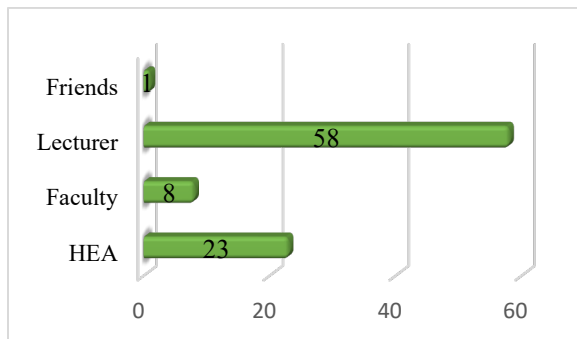


Fig. 5. Sources of information

Table 2. Demographic characteristics of respondents

	Sample	Frequency	Percentage (%)
Gender	Male	36	40
	Female	54	60
Age	18-22	81	90
	23-27	8	8.9
	>27	1	1.1
Educational Level	Diploma	57	63.3
	Bachelor	32	35.6
	Master	1	1.1
Field	Science & Technology	76	84.4
	Science Social	14	15.6

Fig. 5 displays the various sources from which students received information about UnityHub. The data reveals that 58 students learned about UnityHub from their lecturers. Additionally, 23 students received information from the Academic Affairs Division, 8 from faculty, and only 1 student from friends.

As shown in Fig. 6, regarding the purpose of using UnityHub, 46% of students agreed that UnityHub saves time, while 28% found it easy to use. Additionally, 17% of students reported using it at the request of their faculty, and 9% used it out of personal interest. Since UnityHub is easily accessible on mobile devices, students can register for groups anytime, anywhere, and on any platform. Users do not require formal training to use the application, as its simple design allows them to join class groups without navigating multiple web pages.

A similar study by Lee and Kozar (2012) found that perceived ease of use, usefulness, and trust in an application significantly influence user satisfaction. UnityHub users provided positive feedback, recommending that the application be applied or replicated across other faculties. Students responded

favourably to the website's design, navigation, and functionality, which encouraged them to use the application.

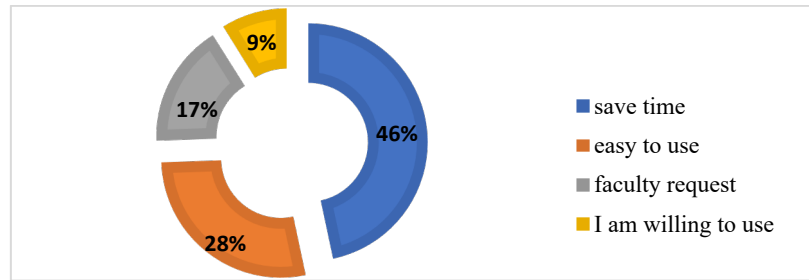


Fig. 6. Purpose in using UnityHub

5.2 Reliability test

Cronbach's Alpha was calculated to assess the reliability of the survey constructs. As shown in Table 3, each Cronbach's Alpha score exceeds 0.80. According to Hinton et al. (2004), Cronbach's Alpha values above 0.50 indicate moderate reliability, while values of 0.70 and above indicate high reliability for dependent and independent variables. The Cronbach's Alpha scores confirm that all measurement items across variables are consistent.

Table 3. Summary of Cronbach's Alpha

Variable	Items	Cronbach's Alpha
Satisfaction	4	0.868
Self-efficacy	4	0.912
Perceived usefulness	5	0.914
Perceived ease of use	4	0.924
Intention to use	3	0.901
Attitude towards use	4	0.872

Table 4. Mean and standard deviation of item

Item	Mean	Standard Deviation
Self-Efficacy		
I am confident of using UnityHub even if there is no one around to show me how to do it.	3.43	0.794
I am confident of using UnityHub even if I have never used such a system before.	3.51	0.811
I am confident of using UnityHub even if I don't have a user manual for reference.	3.38	0.800
I am confident of using UnityHub without assistance from anyone.	3.34	0.837
Satisfaction		
The UnityHub is effective for gathering students' group information.	4.16	0.579
The UnityHub is efficient for time saving.	4.16	0.669
The UnityHub is efficient for cost saving.	4.09	0.647
The UnityHub is efficient for reducing human error.	4.02	0.734
Perceived Usefulness		
Joining class groups is a faster way with UnityHub.	3.82	0.743
UnityHub would enhance my effectiveness in studying.	3.61	0.730
UnityHub would improve the class group joining process.	3.71	0.738
UnityHub facilitates student joining class group.	3.74	0.728
I find UnityHub useful for my studies.	3.67	0.703
Perceived Ease of Use		
The features in UnityHub are clear and understandable.	3.74	0.663
The UnityHub is flexible to access anywhere.	3.80	0.782
The use of UnityHub is very easy to be skilled.	3.77	0.765

I find that UnityHub is easy to use.	3.80	0.767
Intention to Use		
I intend to use UnityHub to enhance communication among lecturers and students	3.77	0.720
I intend to use UnityHub at the beginning of the semester.	3.70	0.741
I intend to use UnityHub every semester.	3.68	0.668
Attitude Toward Use		
Using UnityHub at the beginning of the semester is good.	3.84	0.669
I like the idea of using UnityHub.	3.78	0.700
I have a positive attitude towards using UnityHub.	3.79	0.727
I find the use of UnityHub is a good idea for my study process.	3.79	0.727

Table 4 presents the mean and standard deviation of each item for each construct. All items scored above 3.00, ranging from 3.34 to 4.16, indicating that students generally had moderate to high levels of acceptance of UnityHub. The results encompass various aspects related to UnityHub acceptance, including self-efficacy, satisfaction, perceived usefulness, perceived ease of use, intention to use, and attitude toward the platform.

In terms of self-efficacy, the results indicate that students have a moderate level of confidence in their ability to use UnityHub effectively. On average, they reported being able to navigate the platform without external guidance (mean = 3.43), even without prior experience with similar systems (mean = 3.51) or a user manual (mean = 3.38). Confidence in using the system independently, without assistance from others, was slightly lower (mean = 3.34). These results suggest that while students are reasonably self-reliant, their confidence remains moderate. This aligns with Meng, Qian and Zhang, Qi (2023), who highlighted that self-efficacy is often situational and may require additional training support, and with Mohd Basar, Zulaikha et al. (2021), who observed that learners adapt to new systems even when prior exposure is limited.

Regarding satisfaction, the findings reveal consistently high ratings across all items. Students rated UnityHub as effective for gathering group information (mean = 4.16), efficient for saving time (mean = 4.16), cost-efficient (mean = 4.09), and capable of reducing human error (mean = 4.02). Compared to self-efficacy, which was moderate, satisfaction scores were notably higher, indicating that students perceive strong benefits once they begin using the system. This is consistent with Alyami et al. (2021), who emphasized time management as a critical factor in student performance, and Alhusaini (2023), who also reported that efficiency-driven platforms improve academic experience.

Perceptions of usefulness yielded moderately positive results. Students considered UnityHub helpful in speeding up the process of joining class groups (mean = 3.82) and improving their effectiveness in studying (mean = 3.61). They also rated it as beneficial for enhancing the group-joining process (mean = 3.71), facilitating participation (mean = 3.74), and useful overall (mean = 3.67). Although these scores are positive, they are lower than those for satisfaction, suggesting that while students appreciate UnityHub's role in streamlining processes, they may not yet see it as directly transformative for academic performance. These findings are in line with Hollister et al. (2022), who observed that digital platforms enhance engagement but may vary in perceived academic impact, and Nurkaliza (2018), who noted that usefulness perceptions are shaped by system integration with learning activities.

The construct of ease of use was also rated positively, with mean scores for clarity of features (3.74), accessibility (3.80), and ease of learning (3.77), resulting in an overall mean of 3.80. These results show that students find UnityHub straightforward and convenient to use, a factor that directly supports adoption. This resonates with Bashir (2018), who highlighted the importance of user-friendly design for sustained engagement. Compared with usefulness, ease of use scored slightly higher, suggesting that accessibility and simplicity are stronger drivers of student acceptance.

Students' intention to use UnityHub showed moderate to high results, with mean scores of 3.77 for using it to enhance communication, 3.70 for using it at the start of the semester, and 3.68 for continued use every semester. These scores suggest that students are open to adopting UnityHub as part of their academic

routines, though their intention is slightly lower than their reported satisfaction. This finding reflects Abbas et al. (2019) and Mazana et al. (2019), who both noted that students' adoption of new systems is shaped by both functional efficiency and perceived long-term value.

Finally, students' attitudes toward using UnityHub were favourable. They rated positively the idea of using it at the beginning of the semester (mean = 3.84) and expressed strong liking for the platform (mean = 4.02). Their positive attitude overall (mean = 3.79) and perception that UnityHub is a good idea for supporting their study process (mean = 3.79) reinforce a generally welcoming disposition toward the system. This is consistent with Rawashdeh et al. (2021), who highlighted students' positive attitudes toward e-learning platforms, and Abbas et al. (2019), who found that favourable attitudes strongly influence behavioural intention.

Across constructs, satisfaction recorded the highest mean score (4.16), indicating that efficiency benefits are the strongest driver of student approval. Ease of use and attitude also scored high (around 3.8), reflecting positive experiences with accessibility and favourable overall perceptions. Perceived usefulness and intention to use were moderate (3.7), showing that while students value the system, they may not yet see it as fully essential for academic tasks. Self-efficacy scored lowest (3.42), suggesting that although students can use UnityHub independently, additional support or training could further strengthen confidence. Collectively, these results demonstrate that while UnityHub is well received, future improvements could focus on enhancing its perceived academic usefulness and supporting user self-efficacy.

6. CONCLUSION

This study assessed students' acceptance of UnityHub at Universiti Teknologi MARA (UiTM) across six constructs of the Technology Acceptance Model (TAM). The findings show that satisfaction was rated highest (mean = 4.16), highlighting strong approval of the platform's efficiency in saving time, reducing costs, and minimizing errors. Ease of use (mean = 3.78) and attitude (mean = 3.80) were also high, indicating positive perceptions of accessibility and overall design. By contrast, perceived usefulness (mean = 3.71) and intention to use (mean = 3.72) were moderate, suggesting that while students value UnityHub's role in group management, its academic impact could be strengthened. Self-efficacy was lowest (mean = 3.42), showing that some students may need additional support or training.

These results meet the study objectives by providing a comprehensive overview of acceptance and identifying relative strengths and weaknesses. Practical implications include improving academic features to enhance perceived usefulness and offering onboarding tools to build self-efficacy. Limitations include the single-department sample, which reduces generalizability. Future studies should test UnityHub across multiple faculties and campuses and compare its acceptance with other platforms.

Overall, UnityHub has achieved a generally positive level of acceptance among students, particularly in satisfaction and ease of use, but opportunities remain to strengthen its perceived usefulness and user self-efficacy. Addressing these areas will be crucial to ensuring its sustainable role as an academic communication and management tool in hybrid learning environments.

7. ACKNOWLEDGEMENT/FUNDINGS

The authors wish to express their gratitude to Universiti Teknologi Mara (UiTM), Cawangan Negeri Pulau Pinang, Kampus Permatang Pauh, Malaysia, for their generous support in providing the data and facilities for this research.

8. CONFLICT OF INTEREST STATEMENT

The authors declare that this research was conducted without any personal, commercial, or financial conflicts of interest and affirm there are no conflicts with the funders.

9. AUTHORS' CONTRIBUTION

Norazah Omar: Data Analysis, Results and Discussion and Editing, **Nurhafizah Ahmad:** Methodology, Results and Discussion and Editing, **Jamal Othman:** Literature Review, Conclusion and Abstracts, **Rozita Kadar:** Introduction and System Development.

10. REFERENCES

- Abbas, J., Aman, J., Nurunnabi, M., & Bano, S. (2019). The impact of social media on learning behavior for sustainable education: Evidence from students of selected universities in Pakistan. *Sustainability*, 11(6), 1683. <https://doi.org/10.3390/su11061683>
- Al Husaini, Y., & Ahmad Shukor, N. S. (2023). Factors affecting students' academic performance: A review. *International Journal of Academic Research in Progressive Education and Development*, 12(1), 284–294.
- Alyami, A., Abdulwahed, A., Azhar, A., Binsaddik, A., & Bafaraj, S. (2021). Impact of time-management on the student's academic performance: A cross-sectional study. *Creative Education*, 12(3), 471–485. <https://doi.org/10.4236/ce.2021.123033>
- Ataman, A. (2023, April 8). *6 user acceptance testing best practices in 2023*. AIMultiple. <https://research.aimultiple.com/user-acceptance-testing-best-practices/>
- Bashir, K. (2018). Understanding the role of user interface design in fostering students' learning process in a multimedia courseware learning environment: Insights from a Malaysian case study. *IIUM Journal of Educational Studies*, 5(1), 93–109. <https://doi.org/10.31436/ijes.v5i1.148>
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–339. <https://doi.org/10.2307/249008>
- Gefen, D., Karahanna, E., & Straub, D. W. (2003). Inexperience and experience with online stores: The importance of TAM and trust. *IEEE Transactions on Engineering Management*, 50(3), 307–321. <https://doi.org/10.1109/TEM.2003.817277>
- Ha, S., & Stoel, L. (2009). Consumer e-shopping acceptance: Antecedents in a technology acceptance model. *Journal of Business Research*, 62(5), 565–571. <https://doi.org/10.1016/j.jbusres.2008.06.016>
- Hamilton, T. (2023, April 8). *What is user acceptance testing (UAT)? Examples*. Guru99. <https://www.guru99.com/user-acceptance-testing.html>
- Hollister, B., White, E., Cooley, A., & Salas, S. (2022). Engagement in online learning: Student attitudes and behavior during COVID-19. *Frontiers in Education*, 7, 851019. <https://doi.org/10.3389/feduc.2022.851019>
- Ko, E., Kim, E. Y., & Lee, E. (2009). Modeling consumer adoption of mobile shopping for fashion products in Korea. *Psychology & Marketing*, 26(7), 669–687. <https://doi.org/10.1002/mar.20294>

- Kong, C., Cheung, C., & Yuen, N. (2022). An analysis of the attitudes and behaviours of university students and perceived contextual factors in alternative assessment during the pandemic using the Attitude–Behaviour–Context model. *Heliyon*, 8(10), e11180. <https://doi.org/10.1016/j.heliyon.2022.e11180>
- Lee, Y. K., & Kozar, K. A. (2012). User satisfaction with web-based applications: An exploratory study. *Information & Management*, 49(1), 21–29. <https://doi.org/10.1016/j.im.2011.10.002>
- Mazana, M. Y., Montero, C. S., & Casmir, R. O. (2019). Investigating students' attitude towards learning mathematics. *International Electronic Journal of Mathematics Education*, 14(1), 207–231. <https://doi.org/10.29333/iejme/3997>
- Meng, Q., & Zhang, Q. (2023). The influence of academic self-efficacy on university students' academic performance: The mediating effect of academic engagement. *Sustainability*, 15(7), 5767. <https://doi.org/10.3390/su15075767>
- Mohd Basar, Z., Mansor, A., Jamaludin, K., & Alias, B. (2021). The effectiveness and challenges of online learning for secondary school students: A case study. *Asian Journal of University Education*, 17(3), 119–129. <https://doi.org/10.24191/ajue.v17i3.14514>
- Nugroho, M., Dewanti, P., & Novitasari, B. (2018). The impact of perceived usefulness and perceived ease of use on students' performance in mandatory e-learning use. In *2018 International Conference on Applied Information Technology and Innovation (ICAITI)* (pp. 26–30). IEEE. <https://doi.org/10.1109/ICAITI.2018.8686742>
- Nurkaliza, K. (2018). The role of perceived usefulness and perceived enjoyment in assessing students' intention to use LMS using 3-TUM. [Unpublished manuscript].
- PractiTest. (2023, April 8). *5 user acceptance testing best practices to be more productive*. <https://www.practitest.com/qa-learningcenter/resources/user-acceptance-testing-best-practices/>
- Rawashdeh, A., Youssef, E., Alarab, A., Alarab, M., & Al-Rawashdeh, B. (2021). Advantages and disadvantages of using e-learning in university education: Analyzing students' perspectives. *Electronic Journal of e-Learning*, 19(3), 107–117. <https://doi.org/10.34190/ejel.19.3.2168>
- Smith, J., & Johnson, A. (2018). A systematic literature review to support the selection of user acceptance testing techniques. *ACM Transactions on Software Engineering and Methodology*, 27(3), 1–25. <https://doi.org/10.1145/3183440.3195036>
- Venkatesh, V., Thong, J. Y. L., & Xu, X. (2012). Consumer acceptance and use of information technology: Extending the unified theory of acceptance and use of technology. *MIS Quarterly*, 36(1), 157–178. <https://doi.org/10.2307/41410412>
- Xie, K. L., & Lee, Y. J. (2015). Social media and brand purchase: Quantifying the effects of exposures to earned and owned social media activities in a two-stage decision making model. *Journal of Management Information Systems*, 32(2), 204–238. <https://doi.org/10.1080/07421222.2015.1063297>
- Yuan, G. (2011). Study of implementation of software test management system based on web. In *2011 IEEE 3rd International Conference on Communication Software and Networks (ICCSN)* (pp. 708–711). IEEE. <https://doi.org/10.1109/ICCSN.2011.6014954>

Zaki, A. S., & Ahmad, A. (2017). The level of integration among students at secondary school: A study in Limbang, Sarawak. *The International Journal of Social Sciences and Humanities Invention*, 4(2), 3284–3288. <https://doi.org/10.18535/ijsshi/v4i2.04>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).